

Вопросы распознавания и верификации текстовых документов

В. Л. Арлазаров, О. А. Славин

Федеральный исследовательский центр "Информатика и управление" РАН, Москва, Россия

Аннотация. Работа посвящена концептуальным вопросам анализа текстового содержимого документа. Рассматриваются вопросы использования анализа текстовых документов при создании систем ввода и распознавания деловых бумажных документов. Рассматриваются две основные задачи: определение типа распознаваемого документа и его структуризации. Предложены функции, которые должны обеспечить декомпозицию и решение этих задач. Если принята предлагаемая схема использования текста, то реализация систем ввода может быть различной. То есть, в рамках предложенной концепции анализа документов могут применяться различные алгоритмы распознавания и форматы представления данных.

Ключевые слова: распознавание документов, распознанное слово, сравнение слов.

DOI 10.14357/20718632230306

Введение

Автоматический ввод бумажных документов в компьютер остается важной научной и практической задачей, несмотря на то, что подавляющее большинство документов изначально готовится на компьютере. Этой задаче посвящено множество работ, докладов и даже специальных конференций, таких как:

ICDAR – International Conference on Document Analysis and Recognition (<https://icdar2021.org>, <https://icdar2022.org>, ...);

ICPR – International Conference of Pattern Recognition (<https://icpr2020.net>);

ICMV – International Conference on Machine Vision (<http://icmv.org>).

Главным здесь считается разбиение документа на зоны, в каждой из которых расположен текст, фотография, графическая картинка, являющиеся содержательными частями документа [1-3].

Эти части в исходном или распознанном виде и являются результатом, который выдается в качестве ответа или вводится в базу данных.

Разбиение документа на части осуществляется на основе его графического строения (разделяющих прямых, колонок, абзацев и т.п.) под управлением некоторого описания (шаблона) документа [4, 5]. При этом, привлекаются и элементы распознанных на документе текстов – ключевые слова.

В современных системах документооборота в крупных организациях курсируют документы сотен различных типов, и необходимо уметь определить тип вводимого документа автоматически.

В принципе, задачу определения типа документа можно решить, перебрав имеющиеся шаблоны и попытавшись выбрать шаблон, применение которого дает наилучшее распознавание документа. Но такой подход крайне неэкономичен с точки зрения вычислительных затрат.

В данной работе мы попытаемся разделить задачи идентификации типа и структуризации, опираясь только на лингвистическую составляющую и не привлекая методов сопоставления типа RANSAC, основанных на особых точках.

В большинстве систем распознавания текстов наряду с геометрическим распознаванием изображений символов или групп символов присутствует этап лингвистической обработки [6-9]. На этом этапе проверяется, соответствует ли результат геометрического распознавания лингвистической модели, которой предположительно должен соответствовать результат [9]. В простейшем случае используются проверки слов по словарю.

Если полного словаря возможных слов не существует или он недоступен, часто используются конструкции, основанные на биграммах, триграммах и т.д., когда проверяется сочетаемость пар, троек и т.п. групп букв. Обычно в языковых моделях имеются статистики (вероятности) появления тех или иных сочетаний.

Соответствие результатов геометрического распознавания лингвистической модели служит подтверждением правильности распознавания. Наоборот, их несоответствие требует вмешательства арбитражной системы.

Поскольку система геометрического распознавания выдает результат вместе с оценкой его правильности (слова, его фрагментов или отдельных символов), то арбитражная система может подобрать вариант ответа [10], удовлетворяющий языковой модели и наиболее близкий по некоторой метрике [11, 12] к тому, что выдала геометрическая система распознавания.

1. Распознавание текстов в документе

Ситуация существенно меняется, когда речь идет о распознавании документов. Здесь присутствуют ключевые слова, многие из которых являются заголовками или идентификаторами разделов. С одной стороны, этих слов немного, иногда единицы, и некоторые из них очень «похожи», например, отличаются только окончанием. С другой стороны, значение правильного распознавания таких слов бывает очень велико. Ошибка в них может привести к неправильной трактовке всего документа. Поэтому обычные методы их верификации должны применяться с

большой осторожностью, и мы в данной работе попробуем по-другому подойти к этой задаче.

В обычном тексте лингвистическая система спокойно исправляет, к примеру, орфографическую ошибку, и при этом не имеет значения, была ли она в тексте или возникла в результате распознавания. Однако уже в деловых или исторических документах это далеко не всегда так, а в документах, удостоверяющих личность, почти всегда не так. Ошибка в написании может быть сознательным искажением или, наоборот, результатом неудачной фальсификации. Поэтому организации, осуществляющие ввод идентификационных документов, как правило, негативно относятся к исправлению текста системой распознавания.

Рассмотрим некоторые задачи, в которых присутствует распознавание документов, с точки зрения требований к распознаванию слов.

Если документ представляет собой заполненную форму, то в ней присутствуют как постоянные для форм данного типа поля (поля шаблона), так и заполнение, определяющее собственно экземпляр формы или ее содержание.

Слова шаблона (например, заголовки) определяют его тип и/или идентифицируют зону, в которой будет распознаваться содержание поля. Чаще всего это просто название поля. Словарь слов, составляющих шаблоны, не бывает велик, но имеет ряд особенностей.

Если мы хотим ввести документ, то мы должны распознать некоторые текстовые объекты на нем. Многие тексты можно гораздо быстрее и точнее распознавать, если имеешь представление о том, что это за текст. Возможно, число различных слов в этом тексте не очень велико. Возможно, гарнитура и кегль распознаваемого текста известны. Возможно, мы знаем точно зону, в которой текст надо рассматривать как связный.

Все эти обстоятельства имеют значение, если известен тип распознаваемого документа и его шаблон [13-15]. Таким образом, распознанный тип документа упрощает распознавание текста, а распознанный текст упрощает определение типа.

Возможны и варианты совместной (рекурсивной) работы процессов распознавания текста и определения типа документа. Такое взаимодействие связано, прежде всего, с проведением структуризации документа.

2. Идентификация типа документа

2.1. Особенности идентификации типа документа

Достаточно часто имеет место ситуация, когда тип документа определяется его заголовками, находящимися в верхней части первой страницы.

Другой вариант представлен документами, в которых текстовые зоны отделены друг от друга разграфкой из прямых. Этот случай важен как потому, что встречается часто, так и потому, что разбиение на зоны с помощью разграфки – большое подспорье и при распознавании текстов, и для идентификации типа формы.

Однако если для сканирования используется мобильный телефон или веб-камера, изображение прямоугольной страницы А4 будет подвергнуто не только сдвигу и повороту, как на сканере, но и проективным искажениям. Устранение искажений является отдельной достаточно сложной задачей. Точность, с которой она решается, далеко не всегда достаточна для использования традиционных систем распознавания текстов, впрочем, как и точного определения разделяющих прямых.

2.2. Мешок слов и мешок фрагментов

Таким образом, мы приходим к концепции «мешка» текстовых фрагментов, в частности, мешка слов.

Пусть в шаблоне документа содержится набор слов $W = \{w_i\} 1 \leq i \leq n$. Слова представляют собой некоторый текст. Априори мы не будем предполагать, что тексты между собой не пересекаются.

Для каждого слова w_i имеется описание d_i .

Описание может содержать координаты слова на странице, характеристику шрифта и связи с другими элементами шаблона.

Получив изображение конкретного документа f , система может распознать имеющиеся на документе тексты и сопоставить их со словами того или иного шаблона. Пусть на реальном документе f имеется набор распознанных слов $Wf = \{wf_i\} i \leq N$, и каждое слово имеет описание df_i .

Для верификации типа документа необходимо сопоставить слова шаблона со словами распознанного документа. При распознавании

текста возможны ошибки, одному слову шаблона могут соответствовать несколько мест в распознанном тексте и т.п. Поэтому необходим механизм сопоставления – определения близости двух слов: и текстов, и описаний.

Обычный путь решения подобных вопросов – формирование функции расстояния между двумя объектами [16], в данном случае – словами и их описаниями, а затем, на этой основе, функции расстояния между шаблоном и распознанным документом.

Для реализации такого подхода необходимо построить четыре функции расстояния:

$dist(i, j) = dist(wt_i, wft_j)$ – расстояние между текстами,

$disd(i, j) = disd(wd_i, wfd_j)$ – расстояние между описателями слов,

$disw(i, j) = disw(dist(i, j), disd(i, j))$ – расстояние между словами,

$dish(W, Wf)$ – расстояние между шаблоном и документом.

Мы будем полагать, что значения каждого из расстояний находятся в диапазоне $[0, 1]$.

Эти функции необходимы как при решении задачи определения типа документа, так и для задачи верификации. Однако механизмы их получения и использования различны.

2.3. Идентификация типа документа с помощью шаблонов

Для определения типа документа, минимизируя $dish(W_k, Wf)$ для разных шаблонов W_k , найдется правильный шаблон для данного документа.

Рассмотрим все четыре функции для этого случая чуть подробнее.

$dist(i, j) = dist(wt_i, wft_j)$ обычно определяется бинарно:

$dist(i, j) = \{0 \text{ если } wt_i = wft_j, 1 \text{ в противном случае}\}.$

Это означает, что неправильно распознанные слова просто исключаются из рассмотрения. Расчет идет на то, что в шаблоне содержится достаточно слов, по которым можно идентифицировать тип документа.

Таким образом, получив набор слов в документе F , необходимо найти шаблоны, в которых тексты этих слов содержатся ибо только эти шаблоны, могут претендовать на определение типа документа.

Если $dist(i, j)$ имеет такой простой вид, как указано выше, то, очевидно, можно упростить и функцию $disw$. Именно:

$disw(i, j) = 1$, если $dist(i, j) = 1$

$disw(i, j) = disd(i, j)$, в противном случае.

Т.е. для слов с совпадающими текстами расстояние совпадает с расстоянием между их описателями.

Как правило, описатель слова есть набор признаков, определяющих положение текста слова в документе. Это могут быть координаты слова на документе либо положение слова в некоторой зоне, например, в начале зоны. Для зон, представляющих собой абзац или фрагмент, примыкающий к левому краю страницы, это означает, что текст есть начало строки. Тогда для слов, к которым слева примыкает какой-то текст, этот признак очевидно будет равен 1.

Другой важный признак – обязательность слова в документе. Разумеется, по причинам, рассмотренным выше, этого слова может и не найтись в документе, но признак этот важен при определении функции $dish(W, Wf)$ – расстояния от документа Wf до шаблона W .

Если типов документов много, то сопоставление слов в шаблоне и распознанном документе оставляет место для различных оптимизаций. Прежде всего, это быстрое выделение тех шаблонов, слова из которых есть в тексте документа. Затем проводится экспресс-оценка расстояний между описателями соответствующих слов шаблона и документа и отбраковка шаблонов, которые плохо соответствуют документу. Например, из-за отсутствия в документе большого числа обязательных для данного шаблона слов.

Чем меньше остается вариантов шаблонов, тем больше времени можно уделить верификации шаблонов: разбиению документа на зоны и проверка соответствия шаблону получающейся структуры. Кроме того, после структуризации документа можно оценить взаимное расположение слов в зоне, не укладывающееся в локальное описание слов.

3. Верификация документа

Верификация типа документа и его структуризация, составляющая суть того, что называется «распознаванием документа», ведется по

суущественно другим правилам, хотя, как мы уже отмечали, задачу идентификации типа документа можно свести к его верификации. Для начала это относится к определению функции $dist(i, j)$ – расстоянию между текстами слов в шаблоне и в документе. Поскольку целью здесь является структуризация документа, двузначная функция недостаточна. Необходимо найти все предусмотренные шаблоном зоны и, если какая-то зона определяется единственным обязательным словом, то его нужно найти, даже если оно в документе распознано плохо.

Иногда заранее известно, каким должно быть определенное слово (например, первое слово документа – «Договор», или «Платежное поручение»). В таких случаях возможны следующие варианты совместной работы распознавания и лингвистической системы (ЛС):

Возможны следующие варианты:

а) ЛС подтверждает результат распознавания. Ясно, что в этом случае можно установить $dist(i, j) = 0$,

б) ЛС предлагает вариант достаточно близкий (в смысле метрики) к распознанному. Здесь возможна ошибка распознавания. В этом случае имеем $0 < dist(i, j) < 1$ и значение $dist(i, j)$ должно быть вычислено по некоторой формуле,

в) ЛС предлагает вариант достаточно близкий к распознанному слову, но написано заведомо не оно. Это особенный случай, о котором мы упоминали ранее. С точки зрения структуризации документа можно установить $dist(i, j) = 0$, но заведомое исправление написанного текста рискованно. Возможно, шаблон или зона определены неверно. Нужен арбитраж на уровне $dish(W, Wf)$,

г) близкого по метрике слова в словаре нет. Этот случай может означать, что документ нельзя структуризовать, например, по причине плохого сканирования. Во всех случаях $dist(i, j) = 1$.

В тех случаях, когда слова шаблона не определяют однозначно конкретную зону документа, сопоставление и нахождение зон должно осуществляться по некоторой совокупности слов. Иногда сопоставление может быть обеспечено грубым соответствием группы слов шаблона и документа, текстов и описателей. Но бывает необходимо привлечь и взаимные связи тех или иных слов, которые не могут быть описаны

предложенными здесь функциями. Набор связей между словами достаточно ограничить простейшими отношениями: выше-ниже, слева-справа, дальше по тексту и т.п.

Например, слово «начало работы» должно предшествовать в тексте слову «окончание работы».

Таким образом, для определения какой-то зоны документа должны быть сопоставлены слова шаблона из этой зоны со словами документа, но при этом должны учитываться (проверяться) заданные в шаблоне отношения.

Заключение

В данной статье мы коснулись нескольких вопросов использования текстовой составляющей документа для идентификации его типа и его структуризации. Мы сознательно не пытались предлагать решение многих важных вопросов даже в качестве примера:

- мы не предлагали конкретного языка описания шаблона документа,
- мы не предлагали никаких алгоритмов сопоставления элементов шаблона и текста,
- мы не предлагали никаких форматов (шаблона, результатов структуризации и т.п.),
- мы не пытались даже вкратце описать систему распознавания и ввода документов, в которую должны вписываться обсуждаемые задачи определения типа и структуризации.

Мы всего лишь попытались уяснить, как должны быть связаны проблемы распознавания текстов в документе с его описанием в разных режимах и обратить внимание на некоторые возникающие здесь тонкости.

Мы полагаем, что если предлагаемая схема использования текста принята, то это позволит декомпозировать задачу и сосредоточиться на алгоритмах поиска и перебора шаблонов, разработке функций расстояния и учету связей между элементами шаблона.

Литература

1. Shafait, F., Breuel, T.M.: The Effect of Border Noise on the Performance of Projection-Based Page Segmentation Methods. In *IEEE Transactions on Pattern Analysis and Machine Intelligence* 2011. Vol. 33. No 4. pp. 846-851. (2011). <https://doi.org/10.1109/TPAMI.2010.194>.
2. Melinda, L., Ghanapuram, R., Bhagvati, C.: Document Layout Analysis Using Multigaussian Fitting. *14th IAPR International Conference on Document Analysis and Recognition (ICDAR)* 2017. pp. 747-752. (2017). <https://doi.org/10.1109/ICDAR.2017.127>.
3. Du, X., Wumo, P., Bui, T.D.: Text line segmentation in handwritten documents using Mumford-Shah model. *Pattern Recognition*. Vol. 42. pp. 3136-3145. (2009). <https://doi.org/10.1016/j.patcog.2008.12.021>.
4. Maraj, A., Martin, M.V., Makrehchi, M.: A More Effective Sentence-Wise Text Segmentation Approach Using BERT. In: Lladós, J., Lopresti, D., Uchida, S. (eds) *Document Analysis and Recognition – ICDAR 2021. Lecture Notes in Computer Science*, Springer, Cham. Vol. 12824. (2021). https://doi.org/10.1007/978-3-030-86337-1_16.
5. Mahamoud, I.S., Voerman, J., Coustaty, M., Joseph, A., d'Andecy, V.P., Ogier, J.M. Multimodal Attention-Based Learning for Imbalanced Corporate Documents Classification. In: Lladós, J., Lopresti, D., Uchida, S. (eds) *Document Analysis and Recognition – ICDAR 2021. Lecture Notes in Computer Science*, Springer, Cham. Vol. 12823. (2021). https://doi.org/10.1007/978-3-030-86334-0_15.
6. Nguyen, T.T.H., Jatowt, A., Coustaty, M., Nguyen, N.V., Doucet, A.: Post-OCR error detection by generating plausible candidates. *International Conference on Document Analysis and Recognition, ICDAR 2019*. pp. 876-881. (2019). <https://doi.org/10.1109/ICDAR.2019.00145>.
7. Guo, C.-Y., Tang, Y.Y., Liu, C.-S., Duan, J.: A Japanese OCR post-processing approach based on dictionary matching. *International Conference on Wavelet Analysis and Pattern Recognition*. pp. 22-26. (2013). <https://doi.org/10.1109/ICWAPR.2013.6599286>.
8. Kissos, I., Dershowitz, N.: OCR Error Correction Using Character Correction and Feature-Based Word Classification. *12th IAPR Int. Work. Doc. Anal. Syst. DAS 2016*. pp. 198-203. (2016). <https://doi.org/10.1109/DAS.2016.44>.
9. Bulatov, K., Manzhikov, T., Slavin, O., Faradjev, I., Janiszewski, I.: Trigram-based algorithms for OCR result correction. *Proc. SPIE 10341. Ninth International Conference on Machine Vision (ICMV 2016)*. 103410O, (2017). <https://doi.org/10.1117/12.2268559>.
10. Slavin, O., Janiszewski, I.: Extraction of information fields in administrative documents using constellations of special text points. *Cyber-Physical Systems: Intelligent Models and Algorithms*. Springer Nature Switzerland AG. Vol.417, pp. 267-279. (2023). <https://doi.org/10.1007/978-3-030-95116-0>.
11. Deza, M.M., Deza, E.: *Encyclopedia of distances* // Springer-Verlag, Berlin, 2009, xiv+590 pp.
12. Slavin, O., Andreeva, E., Putincev, D.: Application of modified Levenshtein distance for classification of noisy business document images. *The 14th International Conference on Machine Vision (ICMV 2021)*, November 08-12, 2021. Rome, Italy. *Proceedings Volume 12084, Fourteenth International Conference on Machine Vision (ICMV 2021)*; 120840B (2022) DOI: 10.1117/12.2623437.
13. Postnikov, V.V.: Flexible Forms Identification. *Proceedings of the 5th German-Russian Workshop on Pattern Recognition and Image Understanding (GRWS98)*. Hamburg: Infix, 1999.

14. Postnikov, VV.: Identification and Recognition of Documents with a Predefined Structure // Pattern Recognition and Image Analysis. Vol. 13. No. 2. pp. 332–334. (2003).
15. Марченко, АЕ., Ершов, ЕИ., Гладиллин, С.А.: Система разбора документа, заданного атрибутами структурных элементов и отношениями между структурными элементами. Труды ИСА РАН, Т. 67. No. 4. (2017). С. 87-97.
16. Yujian, L., Bo, L.: A Normalized Levenshtein Distance Metric. IEEE Transactions on Pattern Analysis and Machine Intelligence. Vol. 29, No. 6, pp. 1091-1095 (2007).

Славин Олег Анатольевич. Федеральное государственное учреждение "Федеральный исследовательский центр "Информатика и управление" Российской академии наук", Москва, Россия. Главный научный сотрудник, доктор технических наук, доцент. Область научных интересов: распознавание образов, информационные системы. E-mail: oslavin@isa.ru

Арлазаров Владимир Львович. Федеральное государственное учреждение "Федеральный исследовательский центр "Информатика и управление" Российской академии наук", Москва, Россия. Заведующий отделением, Чл.-корр. РАН, профессор. Область научных интересов: теория графов, распознавания образов, программирование. E-mail: arl@isa.ru

Issues of Recognition and Verification of Text Documents

O. A. Slavin, V. L. Arlazarov

Federal Research Center "Computer Science and Control" of Russian Academy of Sciences, Moscow, Russia

Abstract. The paper deals with the issues of using the analysis of text documents when creating systems for input and recognition of business paper documents. Two main tasks are considered: determining the type of a recognized document and its structuring. Functions are proposed that should provide decomposition and solution of these problems. The work is devoted to the conceptual issues of the analysis of the text content of the document. If the proposed scheme for using text is accepted, then the implementation of input systems may be different. That is, within the framework of the proposed concept of document analysis, various recognition algorithms and data presentation formats can be used.

Keywords: document recognition, recognized word, word comparison.

DOI 10.14357/20718632230306

References

1. Shafait, F., Breuel, TM.: The Effect of Border Noise on the Performance of Projection-Based Page Segmentation Methods. In IEEE Transactions on Pattern Analysis and Machine Intelligence 2011. Vol. 33. No 4. pp. 846-851. (2011). <https://doi.org/10.1109/TPAMI.2010.194>.
2. Melinda, L., Ghanapuram, R., Bhagvati. C.: Document Layout Analysis Using Multigaussian Fitting. 14th IAPR International Conference on Document Analysis and Recognition (ICDAR) 2017. pp. 747-752. (2017). <https://doi.org/10.1109/ICDAR.2017.127>.
3. Du, X., Wumo, P., Bui, TD.: Text line segmentation in handwritten documents using Mumford-Shah model. Pattern Recognition. Vol. 42. pp. 3136-3145. (2009). <https://doi.org/10.1016/j.patcog.2008.12.021>.
4. Maraj, A., Martin, MV., Makrehchi, M.: A More Effective Sentence-Wise Text Segmentation Approach Using BERT. In: Lladós, J., Lopresti, D., Uchida, S. (eds) Document Analysis and Recognition – ICDAR 2021. Lecture Notes in Computer Science, Springer, Cham. Vol. 12824. (2021). https://doi.org/10.1007/978-3-030-86337-1_16
5. Mahamoud, IS., Voerman, J., Coustaty, M., Joseph, A., d'Andecy, VP., Ogier, JM. Multimodal Attention-Based Learning for Imbalanced Corporate Documents Classification. In: Lladós, J., Lopresti, D., Uchida, S. (eds) Document Analysis and Recognition – ICDAR 2021. Lecture Notes in Computer Science, Springer, Cham. Vol. 12823. (2021). https://doi.org/10.1007/978-3-030-86334-0_15
6. Nguyen, TTH., Jatowt, A., Coustaty, M., Nguyen, NV., Doucet, A.: Post-OCR error detection by generating plausible candidates. International Conference on Document Analysis and Recognition, ICDAR 2019. pp. 876–881. (2019). <https://doi.org/10.1109/ICDAR.2019.00145>.
7. Guo, C-Y., Tang, YY., Liu, C-S., Duan, J.: A Japanese OCR post-processing approach based on dictionary matching. International Conference on Wavelet Analysis and Pattern Recognition. pp. 22-26. (2013). <https://doi.org/10.1109/ICWAPR.2013.6599286>.

8. Kissos, I, Dershowitz, N.: OCR Error Correction Using Character Correction and Feature-Based Word Classification. 12th IAPR Int. Work. Doc. Anal. Syst. DAS 2016. pp. 198–203. (2016). <https://doi.org/10.1109/DAS.2016.44>.
9. Bulatov, K., Manzhikov, T., Slavin, O., Faradjev, I., Janiszewski, I.: Trigram-based algorithms for OCR result correction. Proc. SPIE 10341. Ninth International Conference on Machine Vision (ICMV 2016). 103410O, (2017). <https://doi.org/10.1117/12.2268559>.
10. Slavin, O., Janiszewski, I.: Extraction of information fields in administrative documents using constellations of special text points. Cyber-Physical Systems: Intelligent Models and Algorithms. Springer Nature Switzerland AG.. Vol. 417, pp. 267–279. (2023). <https://doi.org/10.1007/978-3-030-95116-0>.
11. Deza, MM., Deza, E.: Encyclopedia of distances // Springer-Verlag, Berlin, 2009, xiv+590 pp.
12. Slavin, O., Andreeva, E., Putincev, D.: Application of modified Levenshtein distance for classification of noisy business document images. The 14th International Conference on Machine Vision (ICMV 2021), November 08-12, 2021. Rome, Italy. Proceedings Volume 12084, Fourteenth International Conference on Machine Vision (ICMV 2021); 120840B (2022) DOI: 10.1117/12.2623437.
13. Postnikov, VV.: Flexible Forms Identification. Proceedings of the 5th German-Russian Workshop on Pattern Recognition and Image Understanding (GRWS98). Hamburg: Infix, 1999.
14. Postnikov, VV.: Identification and Recognition of Documents with a Predefined Structure // Pattern Recognition and Image Analysis. Vol. 13. No. 2. pp. 332–334. (2003).
15. Marchenko A.E., Ershov E.I., Gladilin S.A. Sistema razbora dokumenta, zadannogo atributami strukturnykh elementov i otnosheniyami mezhdru strukturnymi elementami [The system for parsing a document specified by attributes of structural elements and the relations between structural elements]. Trudy ISA RAN. Vol 67, No. 4, pp. 87-97. (2017). (In Russian).
16. Yujian, L., Bo, L.: A Normalized Levenshtein Distance Metric. IEEE Transactions on Pat-tern Analysis and Machine Intelligence. Vol. 29. No. 6, pp. 1091-1095. (2007).

Slavin Oleg A. Federal Research Center “Computer Science and Control” of Russian Academy of Sciences, Moscow, Russia. Lead researcher, Doctor of Technical Sciences, Docent. Scientific interests: artificial intelligence, machine learning, recognition systems, information technology. E-mail: oslavin@isa.ru

Arlazarov Vladimir L. Federal Research Center “Computer Science and Control” of Russian Academy of Sciences, Moscow, Russia. Head of the department, Corresponding member RAS, Professor. Scientific interests: artificial intelligence, recognition systems, information technology, programming. E-mail: arl@isa.ru