

Использование модуля сверточного внимания для интерпретации результатов работы сиамской нейронной сети при идентификации рукописных подписей

 В.А. Мищук

ФГАОУ ВО «Российский университет дружбы народов имени Патриса Лумумбы», Москва 117198, Россия

Аннотация. В настоящей работе рассматривается возможность применения сиамской сверточной нейронной сети (SNN) с интегрированным модулем сверточного внимания (CBAM) в идентификационных исследованиях рукописных подписей. Одним из факторов, тормозящих процесс внедрения искусственных нейронных сетей в процесс производства судебно-экспертных исследований, является их низкая степень интерпретируемости. Из-за этого исследователю сложно определить, какие именно закономерности были выявлены нейросетевым алгоритмом и какие из них легли в основу полученного прогноза. Кроме того, в большинстве современных работ, посвященных анализу почерка, специалисты используют «классический» подход к определению авторства рукописи, при котором эта задача рассматривается как частный случай классификации. Однако данный способ часто приводит к ошибкам II рода, из-за чего, на взгляд авторов, использование классификационных алгоритмов для решения идентификационных задач неприемлемо. Вместо этого авторы предлагают обратить внимание на архитектуру SNN. Для подтверждения этих тезисов в рамках настоящей работы были проведены эксперименты, в ходе которых удалось установить, что современные механизмы внимания, в частности модуль CBAM, способны частично интерпретировать полученные нейросетью результаты. Применение SNN, в свою очередь, позволяет минимизировать число ошибок II рода по сравнению с «классической» классификационной системой.

Ключевые слова: идентификация подписи, интерпретируемость, искусственные нейронные сети, механизм внимания, сиамская нейронная сеть, судебно-почерковедческая экспертиза, Convolutional Block Attention Module (CBAM)

Для цитирования: Мищук В.А. Использование модуля сверточного внимания для интерпретации результатов работы сиамской нейронной сети при идентификации рукописных подписей // Теория и практика судебной экспертизы. 2025. Т. 20. № 2. С. 65–81.

<https://doi.org/10.30764/1819-2785-2025-2-65-81>

Using Convolutional Block Attention Module for Interpreting the Results of Artificial Neural Networks Operation in Handwritten Signature Identification

 Vsevolod A. Mishchuk

Peoples' Friendship University of Russia named after Patrice Lumumba, Moscow 117198, Russia

Abstract. This study explores potential application of a Siamese Convolutional Neural Network (SNN) integrated with a Convolutional Block Attention Module (CBAM) in the field of handwritten signature identification. One well-known obstacle to the introduction of artificial neural networks into forensic expert practice is the low level of their interpretability.

This limitation makes it difficult for experts to determine what patterns were identified by the neural network algorithm, and which ones of them form the basis of its forecast. Furthermore, in most contemporary studies on handwriting analysis, specialists tend to use the “classical” authorship attribution approach, treating the task as a particular case of classification.

However, in authors' view, this method often leads to Type II errors making the use of classification algorithms for identification purposes unacceptable. As an alternative, the authors propose to focus on the SNN architecture. To support these claims, a series of experiments were conducted as part of this study, demonstrating that modern mechanisms of attention – particularly CBAM – can partially interpret

the neural network results. Additionally, the use of SNN helps to minimize the number of Type II errors compared to the traditional classification-based approach.

Keywords: *signature identification, interpretability, artificial neural networks, attention mechanism, siamese neural network, forensic handwriting examination, Convolutional Block Attention Module (CBAM)*

For citation: Mishchuk V.A. Using Convolutional Block Attention Module for Interpreting the Results of Artificial Neural Networks Operation in Handwritten Signature Identification. *Theory and Practice of Forensic Science*. 2025. Vol. 20. No. 2. P. 65–81. (In Russ.). <https://doi.org/10.30764/1819-2785-2025-2-65-81>

Введение

В настоящее время одной из актуальных для юридического сообщества тем является интеграция искусственных нейронных сетей (ИНС) в судебно-экспертную деятельность (СЭД). К примеру, эксперты-криминалисты системы МВД России с недавнего времени стали использовать в своей практической деятельности АДИС¹ «Папилон-9», в которой ИНС применяется для автоматического кодирования отсканированных изображений следов рук с последующим поиском и сравнением с другими следами, хранящимися в базе данных². Также имеются работы, в которых нейронные сети используются для решения задач, связанных с судебной баллистикой [1, 2].

В то же время, как верно замечают многие исследователи, внедрять подобные разработки в практику производства именно судебно-экспертных исследований разного рода/вида нужно с большой осторожностью, поскольку этот процесс связан с многочисленными проблемами методического, организационного и правового характера [3–5]. Перечисленные выше системы, несмотря на их тесную связь с судебной экспертизой, фактически представляют собой пример использования специальных знаний в непроцессуальном порядке для получения ориентирующей информации, способствующей расследованию и раскрытию преступлений. Этот процесс скорее относится к криминалистике в широком ее проявлении и оперативно-розыскной деятельности.

АДИС «Папилон» применяется лишь как средство ведения дактилоскопического учета и поиска лиц, вероятно совершивших преступление, по обнаруженным на месте

происшествия следам рук. Результаты подобных проверок, равно как и иные схожие непроцессуальные формы использования специальных знаний, в настоящее время не признаются действующим законодательством источником доказательственной информации по делу. В связи с этим на большинство подобных систем обычно не распространяются те требования, которые предъявляются к заключению судебного эксперта – одной из форм процессуального использования специальных знаний, являющейся доказательством во всех видах судопроизводства. Так, в соответствии со ст. 8 Федерального закона от 31.05.2001 № 73-ФЗ «О государственной судебно-экспертной деятельности в Российской Федерации» результаты экспертного исследования должны основываться: «...на положениях, дающих возможность проверить обоснованность и достоверность сделанных выводов на базе общепринятых научных и практических данных». Кроме того, применяемые в экспертизе методы должны отвечать критериям научной обоснованности, точности и надежности получаемых результатов [6], а само заключение должно отражать: «... содержание и результаты исследований с указанием примененных методов»³. Иными словами, в любой конкретной судебной экспертизе получаемые результаты и выводы, основанные на них, должны быть интерпретируемыми для возможности их оценки как другими экспертами/специалистами, так и иными участниками процесса, которые не являются сведущими лицами.

Нейронные сети, как отмечают многие ученые, относятся к алгоритмическим системам типа «черный ящик». Это означает, что исследователь или разработчик даже при знании того, как в целом происходит процесс обработки входных данных, не

¹ Автоматизированная дактилоскопическая информационная система.

² Папилон-Нейро – +8% identifications in AFIS (AFIS) // Системы Папилон. <https://www.papillon.ru/products/support/neuro/?ysclid=lxbw9lsxeu279905013>

³ Ст. 25 Федерального закона от 31.05.2001 № 73-ФЗ «О государственной судебно-экспертной деятельности в Российской Федерации».

может с точностью сказать, какие именно закономерности в этих данных выявляет система и какие из них использует для решения поставленной перед ней задачи. Данный факт существенно ограничивает возможности использования нейросетевых алгоритмов для решения задач того или иного рода/вида судебной экспертизы и отрицательно сказывается на воспроизводимости и точности получаемых результатов. Известны случаи, когда подобная система для расчета своего «вывода» использовала не сами исследуемые данные, а случайные шумы, сторонние объекты на изображении, его отдельные особенности в виде размера, метаданных и т. п. [7, 8].

Особенно остро вопрос интерпретируемости получаемых с помощью ИНС прогнозов стоит с судебно-почерковедческой экспертизе, где качественно-описательная экспертная методика исследования различных рукописных объектов является доминирующей в силу специфики и сложности человеческого почерка. Разумеется, в практике этого рода экспертизы используются количественные и комплексные методы и методики. Более того, как показывает исторический анализ, отечественные криминалисты еще в 1960-х гг. обратили внимание на первые прототипы современных ИНС – перцептроны – и уже тогда изучали возможность их применения для решения отдельных задач судебно-почерковедческой экспертизы. Однако, как отмечают В.В. Устинов [9] и В.А. Мещеряков [10], несмотря на наличие подобных разработок и методик, практикующие эксперты используют их достаточно редко из-за недоверия к получаемым результатам, трудоемкости и в целом недостаточного уровня подготовки к работе с этими методиками.

Следует также затронуть проблему решения идентификационных задач⁴ дихотомическими алгоритмами, на которую впервые обратил внимание Л.Г. Эджубов в 1970-х гг. [11]. Суть данной проблемы заключается в том, что нейросетевая и иная схожая система более чувствительна к ошибкам II рода, когда, например, среди лиц, которые предположительно могли выполнить исследуемую рукопись, нет ее истинного исполнителя. В то же время ИНС из-за особенностей своей архитектуры и

механизмов принятия решений с большей долей вероятности отнесет эту рукопись к автору с наиболее схожим почерком.

Исходя из вышеизложенного, нами была предпринята попытка по решению этих проблем посредством использования сиамской сверточной нейронной сети с интегрированным в ее архитектуру механизмом внимания (attention). Был выбран модуль сверточного внимания (Convolutional Block Attention Module – CBAM) как наиболее хорошо адаптированный для работы со сверточными нейронными сетями (Convolutional neural network – CNN) и использующий сравнительно небольшое количество вычислительных ресурсов компьютера. По задумке модуль CBAM должен выполнять роль интерпретатора, с помощью которого планировалось отображать те области исследуемой рукописи, в которых предположительно содержались наиболее важные идентификационные признаки почерка. В свою очередь особенности функционирования сиамской ИНС (Siamese Neural Network – SNN) позволили минимизировать количество ошибок II рода, возникающих в условиях ограниченного набора данных в виде предоставленных образцов почерка небольшого круга известных суду и следствию лиц.

1. Краткий анализ предыдущих работ

В своем исследовании мы, в первую очередь, ориентировались на работу специалистов из Республики Корея, которые разработали модуль CBAM [12], а также на публикацию ученых из Пекинского университета Цзяотун, в которой описано применение двухпутевой CNN с CBAM для дифференциации писателей по половому признаку путем анализа рукописных текстов [13]. Тем не менее, кратко рассмотрим и другие разработки, нацеленные на интерпретацию результатов работы ИНС при анализе почерка. Например, экспериментальный программный комплекс «Фрося», предложенный М.В. Бобовкиным, О.А. Диденко и А.Е. Нестеровым [14], по задумке авторов должен при помощи нейронных сетей собирать статистические данные о частных признаках почерка и выводить интерактивные подсказки по дифференциации того или иного признака. Иными словами, система анализирует каждый отдельный символ в рукописи и выделяет определенные частные признаки. Так, авторам удалось обучить

⁴ Здесь и далее термин «идентификация» используется в его криминалистическом смысле – установление индивидуально-конкретного тождества объекта.

нейросеть находить и дифференцировать форму движений при выполнении и соединении элементов в некоторых буквах.

Схожая идея прослеживается в работе специалистов Института автоматизации Китайской академии наук, которые использовали нейросеть для изучения рукописей, выполненных при помощи устройств электронного ввода – стилуса и графического планшета [15]. Для поиска признаков почерка в таком типе данных специалисты применяли двунаправленную сеть долгой краткосрочной памяти (Long short-term memory – LSTM), а не CNN, как в комплексе «Фрося». Это обусловлено тем, что подобные «цифровые» рукописи представляют собой последовательность значений положения, силы давления, углов выполнения и иных характеристик пера стилуса, которые фиксируются электронным устройством. При этом сами признаки представляют собой отдельные участки этой последовательности. Как показывает современная практика, лучше всего с задачей изучения последовательностей справляются рекуррентные ИНС (Recurrent neural network – RNN), LSTM, архитектуры типа Transformer и схожие модели. Подобный подход в совокупности с ансамблевым методом анализа, который подробнее описывается ниже, позволяет определять участки, которые с наибольшей вероятностью характеризуют почерк конкретного автора.

Важный вклад в анализ рукописей, выполненных традиционным способом, внесли Хсин-Хсиунг Као и Че-Йен Вен [16]. Они использовали метод визуализации карт значимости, которые предложили К. Симонян, А. Ведалди и Э. Зиссерман. Данный способ интерпретации строится на расчете градиента⁵ входных данных – анализируемого изображения – относительно полученного сетью результата. Как указывают авторы: «... величина производной показывает, какие пиксели должны быть изменены в наименьшей степени, чтобы повлиять на оценку класса в наибольшей степени» [17]. Иными словами, происходит визуализация тех областей объекта, изменение которых приведет к изменению прогноза. Используя этот метод, Као и Вен смогли продемонстрировать, что CNN после ее обучения действительно использует для про-

гнозирования закономерности почерка, а не сторонний шум изображения. Например, они установили, что: «... поворотная точка и пересечение штрихов часто используются в качестве важной основы для проверки подписи» [16].

Помимо визуализации карт значимости существует еще один схожий способ объяснения того, какие закономерности выявляет и использует нейросеть при принятии решения – вычисление «взвешенного значения по градиенту отображения активации класса» (Gradient-weighted Class Activation Mapping – GradCAM) [18]. Если при визуализации карт значимости градиент рассчитывается относительно входного изображения, то в GradCAM это происходит относительно карт признаков, получаемых при анализе этого изображения сверточными слоями сети. Данный метод был использован исследователями из Университета им. Бен-Гуриона для отображения признаков, которые, как предполагается, позволяют дифференцировать документы, написанные в разные периоды истории [19].

Иной подход к вопросу определения, какие признаки использует ИНС при анализе почерка, был предложен М. Марциновским [20], который в качестве основы для собственной модели использовал архитектуру нейросети VGG16, где базовый классификатор был заменен на два полносвязных слоя нейронов. Первый слой, принимающий информацию из сверточных слоев VGG16, состоит из 84 нейронов с сигмоидальной функцией активации. Каждый из них соответствует определенному признаку почерка по системе, предложенной Р. Хубером и А. Хедриком [21]. Предполагается, что нейрон будет выдавать значение «0» или «1», тем самым сигнализируя об отсутствии либо наличии соответствующего признака в исследуемой рукописи. Таким образом, итоговый вектор из этой последовательности «0» и «1» передается в выходной полносвязный слой, состоящий из N-нейронов, где N – число предполагаемых авторов, образцы почерка которых были предоставлены системе для обучения. Как итог, прогнозом ИНС является: «... вектор нулей с единственной «1», которая соответствует рассматриваемому автору (одна метка – идентификатор; множество классов – писатели)» [20]. Таким образом, идентификатор в виде «1» указывает на то проверяемое лицо, которое скорее всего является исполнителем исследуемой рукописи.

⁵ Градиент – это вектор, состоящий из частных производных функции потерь (числовой показатель расхождения между прогнозом и реальными данными), рассчитанных по всем весам синапсов (связи между нейронами) в нейросети.

Как уверяет автор, модель является максимально интерпретируемой, что было подтверждено при помощи метода визуализации карт значимости нейросети, описанного ранее [16]. В то же время «...можно утверждать, что модель выучила умеренное количество неизвестных признаков, которые действительно совпадают по значению и частоте с характеристиками, которые мы определили. Эти признаки могут соответствовать характеристикам или крайней мере сильно коррелировать с ними» [20]. Иными словами, М. Марциновский не дает полной гарантии того, что модель использует именно те признаки почерка, которые были заданы исследователем.

Говоря о решении задачи криминалистической идентификации при помощи ИНС, отметим, что в последнее время в этом направлении чаще всего используют сиамские нейросети. Например, Дж. Бромли, И. Гийон, Я. ЛеКун, Э. Сикингер и Р. Шах, которые одними из первых предложили подобный подход, использовали SNN для идентификации исполнителя рукописи, выполненной на графическом планшете при помощи стилуса [22]. Достигнутая по итогу обучения точность составила порядка 95,5 %.

Из современных исследований можно отметить работу Центра компьютерного зрения Автономного университета Барселоны [23], в которой нейросеть SigNet в среднем достигла точности порядка 85,5 % на различных наборах данных. Нельзя не упомянуть и проект NSP-SigVer, в котором исследователи кафедры криминалистики Уральского государственного юридического университета создали собственную экспериментальную модель SNN, а также собрали большую базу данных подписей, выполненных кириллицей [24]. Также SNN была применена компанией «Т-информ» при разработке АСПИАД⁶ – системы, предназначенной для: «... анализа рукописных материалов на предмет определения неоднородности почерка и принадлежности текста разным авторам»⁷.

Еще один способ решения идентификационных задач связан с использованием байесовских нейронных сетей и ансамблей нейросетей. Оба подхода основаны на предположении, что применение нескольких нейросетей или механизма вероятности

даст вероятностно-статистический результат. Это означает, что если исследуемый образец действительно выполнен писателем, которого нейросеть знает (поскольку она обучалась на его образцах), то она с большей долей вероятности отнесет образец именно к этому автору. Если же анализируемый объект не принадлежит ни к одному из известных нейросети классов, то вероятность его отнесения к любому из них, даже внешне схожему, будет низкой, что можно сравнить с отрицательным экспертным выводом или НПВ⁸. Так, П. Чопра продемонстрировал, что сети подобной архитектуры неплохо справляются с классификацией рукописных цифр, при этом отдельно выделяя те, которые сложно однозначно отнести к какой-то определенной группе [25]⁹. Упомянутые выше специалисты Института автоматизации Китайской академии наук использовали более простой подход в виде ансамблевого метода [15]. В частности, все исходные данные были поделены на фрагменты, для каждого из которых RNN делала прогноз, после чего все полученные результаты усреднялись, формируя итоговый вывод.

Разумеется, это не полный перечень исследований, посвященных интерпретации результатов работы ИНС при анализе почерка, но на наш взгляд, приведенные публикации максимально полно представляют основные направления и подходы в решении этих проблем.

2. Описание архитектуры искусственной нейронной сети

2.1. СВАМ – принцип работы

Для начала обратимся к принципу работы модуля СВАМ, который является результатом развития идеи механизма внимания (attention), предложенного Д. Богдановым, Кенхен Чо и Й. Бенжио для увеличения точности нейросетевых моделей [26]. Механизм внимания в ИНС глобально представляет собой последовательность математических операций, посредством которых можно выделить наиболее важные для составления верного прогноза участки входных данных.

⁶ Аналитическая система почерковедческой идентификации авторов документов.

⁷ Аналитическая система АСПИАД // Т-информ. <https://tinform.ru/solutions/aspiad/>

⁸ Решить поставленный вопрос не представилось возможным.

⁹ Однако обучать байесовские нейронные сложнее, и они не всегда отличаются высокой надежностью.



Рис. 1. Общее устройство модуля СВМ (верх), устройство блока канального внимания (центр), устройство блока пространственного внимания (низ) [12]

Fig. 1. General design of CBAM module (upper), Channel Attention Block Design (middle), Spatial Attention Block Design (lower) [12]

Механизм внимания в СВМ реализован при помощи двух вычислительных блоков – канального и пространственного внимания (рис. 1, верх). Блок канального внимания анализирует карты признаков, сформированные сверточными фильтрами конкретного слоя CNN при анализе изображений, и определяет в них внутриканальные связи через изучение информации о яркости пикселей (рис. 1, центр). Иными словами, при помощи блока канального внимания можно получить информацию о том, какие участки изображения предположительно наиболее важны при прогнозе.

В свою очередь блок пространственного внимания, используя те же карты признаков и информацию, полученную с блока канального внимания, определяет, где именно находятся выявленные важные участки изображения (рис. 1, низ). Таким образом, получаем карты внимания, в которых отражается информация о том, какие признаки являются ключевыми при анализе того или иного объекта и где они расположены.

2.2. Сиамская нейронная сеть как механизм принятия решения при идентификации

Частично мы уже рассмотрели «классический» подход распознавания образов, который применяется многими исследо-

вателями для определения авторства рукописи с помощью ИНС. Так, в его рамках задача идентификации рассматривается как частный случай многоклассовой классификации, когда необходимо отнести исследуемый объект к заранее определенному классу путем создания набора обучающих данных – изображений рукописей потенциальных авторов и целевых меток, указывающих, к какому классу/автору относится каждый из представленных образцов. Затем моделируется сама нейронная сеть, в которой выходной слой обычно содержит N-нейронов, где N – количество классов/авторов в обучающей выборке (рис. 2). После обучения на этом ограниченном наборе данных системой достигается высокая точность и создается впечатление, что она способна успешно идентифицировать образцы.

Однако на практике ситуация выглядит иначе. Когда обученная нейросеть анализирует рукопись, не выполненную ни одним из проверяемых лиц, на образцах почерка которых проводилось обучение системы, она относит ее к исполнителю с наиболее похожим почерком, вместо того чтобы признать, что автором рукописи является некое другое лицо. Соответственно, система ограничена данными, на которых она обучалась, и поэтому более склонна к ошибке II рода.

Здесь будет уместно процитировать Л.Г. Эджубова, который еще в 1970-х гг. указывал следующее: «Алгоритмы подобного класса предназначены для решения задачи дихотомии, то есть разделения множества объектов на два (и более – *прим.*) класса. Нельзя не заметить, что задача дихотомии не совпадает с задачей, стоящей обычно перед экспертом при идентификации» [11]. Разумеется, можно, например, ввести еще один класс, в котором будут представлены изображения рукописей множества различных авторов. Однако следует учесть, что индивидуальность почерка каждого человека не исключает его возможного сходства с почерками других людей. Вполне вероятна ситуация, в которой почерк настоящего и предполагаемого исполнителя будут очень

похожи друг на друга, из-за чего в процессе классификации ИНС все равно будет допускать ошибки, особенно в случаях, когда условный класс «подделок» был сформирован плохо и сильно внешне отличается от анализируемой рукописи.

SNN, в свою очередь, решает данную задачу через определение степени сходства двух пар исследуемых образов данных при использовании двух нейронных сетей с общими обучаемыми параметрами (рис. 3).

Идея и сам процесс обработки данных посредством SNN состоит в следующем: два изображения подаются на нейронные сети, которые преобразуют их в N-мерный вектор. Предполагается, что если анализируемая пара объектов относится к одному классу, то полученные вектора должны об-

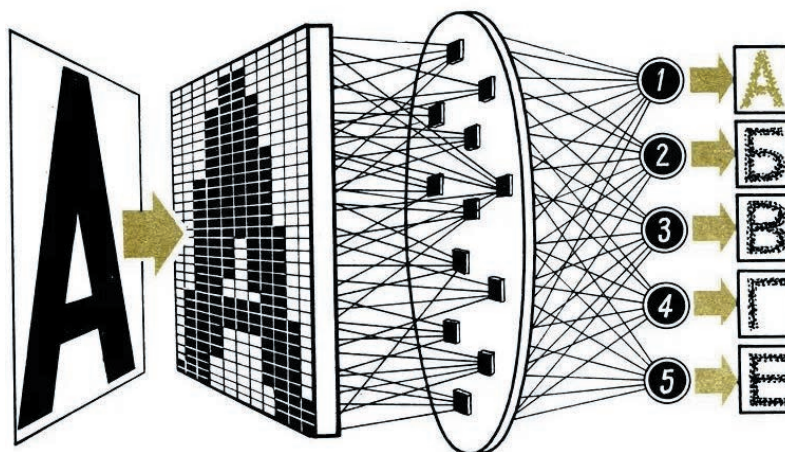


Рис. 2. Пример решения задачи многоклассовой классификации при помощи искусственной нейронной сети. Изображение взято из открытых источников.

Fig. 2. Example of solving a multiclass classification task using an artificial neural network. Image taken from publicly available sources

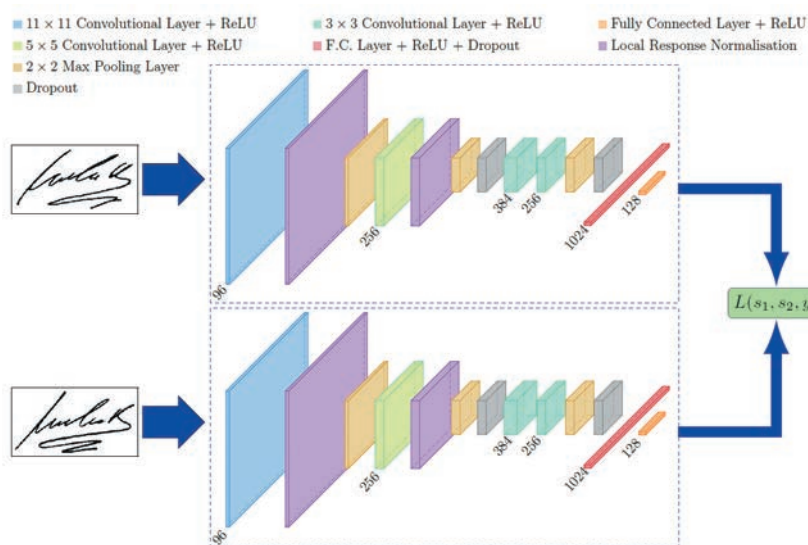


Рис. 3. Устройство и схема работы сиамской нейронной сети [23]

Fig. 3. Design and operation scheme of Siamese neural network [23]

ладать какими-то общими характеристиками, например, находиться друг от друга на минимальном расстоянии в N-мерном пространстве, если использовать такой критерий сходства как евклидово расстояние. Напротив, если объекты не принадлежат к одному и тому же классу, то получаемые SNN вектора должны сильно отличаться, то есть, как в случае с евклидовым расстоянием, находиться как можно дальше друг от друга. При этом само обучение подобной системы происходит примерно так же, как и для обычных ИНС, – при помощи метода обратного распространения ошибки.

2.3. Описание собственной архитектуры ИНС

На основе описанных выше исследований и разработок нами была спроектирована искусственная нейронная сеть, основой для которой послужила CNN ResNet (Residual Network). Особенность этой ИНС заключается в использовании так называемых остаточных блоков (residual block) (рис. 4), которые позволяют легко оптимизировать нейросеть вне зависимости от ее «глубины» (количества слоев), повышая качество обучения системы и минимизируя эффект «затухающих градиентов». После остаточных блоков мы внедрили модули СВМ, частично повторив подход, предложенный исследователями из Пекинского университета Цзяотун [13]. При этом рассматриваемая ИНС имеет следующие существенные особенности и модификации.

1. Модули СВМ не передают информацию с определенного уровня нейросети последующим слоям, чтобы выявляемые блоком attention карты внимания не влияли на результаты работы этих слоев нейросети. Это позволяет извлекать максимум информации из изображения, выявляя как «глубинные» признаки объекта, так и «высокоуровневые», то есть наиболее общие¹⁰. Фактически, блоки attention служат «хранилищем» наиболее важных признаков, полученных с разных уровней базовой нейросети, благодаря чему входное изображение анализируется практически так же, как и в обычной ResNet.

¹⁰ Важно понимать, что в контексте обучения ИНС под признаком понимается не какое-то отдельное свойство объекта, например форма движения при выполнении определенного элемента, а в целом любая потенциально важная информация об этом объекте исследования. Чаще всего такая информация имеет количественный характер.

2. Поскольку полученные карты внимания различаются по размеру, из-за того, что были сформированы на основе признаков, извлеченных с разных уровней ResNet, они интерполируются до размера наибольшей карты. Это позволяет в дальнейшем производить их более точное масштабирование до размеров исходного изображения рукописи. Таким образом, имеется возможность накладывать выявленные карты внимания на изображение с минимальными искажениями, благодаря чему потенциально можно визуально определить наиболее информативные участки рукописи. После интерполяции все полученные карты объединяются в один массив, который обрабатывается финальным модулем СВМ для определения наиболее значимых признаков. В результате условный «классификатор» должен вычислять прогноз на основе наиболее значимой информации о выявленных признаках почерка, а не всего массива возможных данных.

3. Поскольку мы работаем с сиамской ИНС, нам не требуется «классический» классификатор, где количество нейронов выходного слоя соответствует числу предполагаемых классов. Вместо этого используется полносвязный слой из 128 нейронов в качестве выхода ИНС, что равносильно обычному вектору, который описывает сложные входные данные.

В итоге весь процесс обработки данных в приведенной экспериментальной модели проходит следующие этапы (рис. 4):

1. Два изображения рукописи подаются на входные уровни сиамской нейронной сети, обрабатываются входным сверточным слоем.

2. Происходит анализ остаточными сверточными блоками, информация с которых передается в модули СВМ для вычисления карт внимания.

3. Карты внимания интерполируются и объединяются, обрабатываются финальным модулем СВМ.

4. Полученный массив данных передается в Adaptive AvgPool (адаптивное среднее объединение) для снижения размерности.

5. В полносвязных слоях вычисляются 128-мерные векторы для двух изображений. Сравнение векторов происходит по косинусному сходству с пороговым значением 0,75. Если результат больше этого значения, значит, две рукописи были выполнены одним автором, в противном случае – разными.

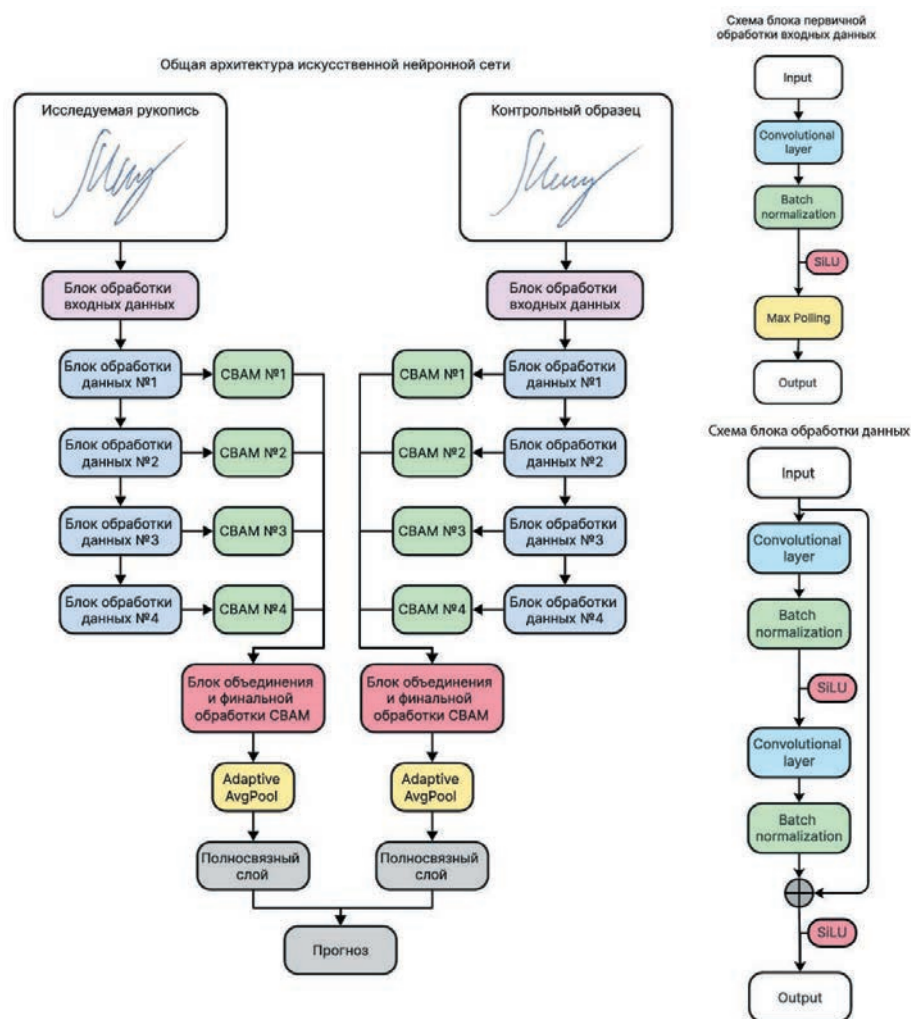


Рис. 4. Архитектура искусственной нейронной сети, использованной в данном исследовании
Fig. 4. The architecture of the artificial neural network used in this study

3. Экспериментальные исследования

3.1. Обучающие данные и постановка задачи

В настоящее время методология судебно-почерковедческой экспертизы позволяет решать широкий спектр задач, однако на практике наиболее типичной из них является идентификация исполнителя рукописной подписи. Именно ее и старались решить в рамках экспериментального исследования.

Идентификация человека по его подписи и почерку в целом является областью интереса не только судебной экспертизы, но и биометрии. Данное научное направление тесно связано с машинным обучением и нейронными сетями, поскольку последние предоставляют эффективные способы обработки и анализа биометрических параметров человека. Сегодня существует большое количество баз данных (dataset),

состоящих из отсканированных изображений рукописей, в том числе и подписей, выполненных различными людьми. Среди основных можно выделить:

- **CEDAR**. Один из наиболее популярных и старых наборов данных, где собраны подписи 55 человек. На каждого из них приходится по 24 образца сфальсифицированной подписи, что позволяет использовать набор для обучения ИНС распознавать настоящие подписи и отличать их от подделок.

- **GPDS**. Более крупный набор данных, содержащий 150–300 подписей от разных людей. Как и в CEDAR, у каждого человека есть 24 настоящих подписи и 30 подделок.

- **ICDAR 2011 SigVer**. В данном популярном наборе данных помимо обычных рукописных есть подписи, выполненные на компьютере с помощью стилуса и планшета, а также подписи, написанные на латинице и иероглифами.

Несмотря на большой объем информации, у перечисленных наборов данных есть общий недостаток – все они сформированы для того, чтобы обучать модели решать задачу верификации исполнителя подписи. Верификацию в сфере компьютерных наук и биометрии можно определить как процесс сопоставления данных между собой. В свою очередь в криминалистике и судебной экспертиологии решается задача идентификации, то есть установления индивидуально-конкретного тождества объекта исследования, в данном случае – установления автора исследуемой подписи. Отличие между задачами состоит в том, что при верификации достаточно установить, выполнена ли рукопись тем человеком, от имени которого она значится, а при идентификации в целом решается вопрос, кто является истинным исполнителем исследуемой рукописи, даже если она выполнена с подражанием почерку другого человека.

К настоящему времени еще не созданы базы данных, при помощи которых можно было бы обучить ИНС решать идентификационные задачи. В связи с этим нами был сформирован собственный небольшой набор данных, который позволил обучить нейронную сеть не только верифицировать человека по подписи, но и проводить идентификацию в криминалистическом смысле. Он состоял из образов подписей-оригиналов в количестве около 100 штук и подписей-имитаций, выполненных шестью людьми также в количестве около 100 штук от каждого исполнителя, и поделен внутри на два поднабора. Первый использовался как для обучения модели, так и для тестирования и включал в себя образцы подписи-оригинала, а также подделки, выполненные четырьмя людьми (обозначен в таблице 1 как

«Own dataset № 1»). Второй использовался только как dataset для тестирования и также состоял из образцов-оригиналов и сфальсифицированных подписей, выполненных двумя оставшимися исполнителями (обозначен в таблице 1 как «Own dataset № 2»).

3.2. Процесс обучение и тестирования ИНС

1. Из вышеперечисленных наборов данных выбирали один, на котором, собственно, и происходило обучение ИНС.

2. Для каждого образца оригинальной подписи случайным образом (с вероятностью 50 %) подбирали либо другой образец-оригинал, либо сфальсифицированную подпись, чтобы сбалансировать обучающий набор данных, тем самым повысив качество обучения.

3. ИНС проверяли на тестовой выборке данных из того же dataset, на котором она была обучена (метка *test*), а также на данных из других наборов, которые были указаны ранее.

В качестве метрики качества обучения была выбрана точность (метка *acc*), которая определялась как отношение количества правильных прогнозов к их общему числу. Кроме того, в таблице 1 отражены значения функции потерь (*loss*) при проведении тестирования и достигнутая моделью точность на конкретном dataset на обучающей выборке (метка *train*).

Поскольку архитектура сиамской нейросети допускает использование одной ИНС, чтобы повысить точность прогнозирования первоначально модели обучали по «классической» схеме, которая предполагает классификацию данных. В частности, модели определяли, каким именно исполнителем была выполнена подпись – вне зависимости

Таблица 1. Результаты обучения искусственной нейронной сети
Table 1. Results of training the artificial neural network

Тестовые данные	CEDAR	GPDS150	ICDAR 2011 SigVer	Own dataset № 1	Own dataset № 2
Обучающие данные	train: acc – 1,0, loss – 0,01488	acc – 0,526	acc – 0,558	acc – 0,711	acc – 0,781
	test: acc – 0,997 loss – 0,00937	loss – 0,50343	loss – 0,46641	loss – 0,31680	loss – 0,31191
CEDAR	acc – 0,690	train: acc – 0,795 loss – 0,23338	acc – 0,621	acc – 0,675	acc – 0,776
	loss – 0,36570	test: acc – 0,746 loss – 0,26843	loss – 0,42439	loss – 0,42810	loss – 0,38037
GPDS150	acc – 0,869	acc – 0,6	train: acc – 0,880 loss – 0,14199	acc – 0,720	acc – 0,801
	loss – 0,17364	loss – 0,48598	test: acc – 0,945 loss – 0,11715	loss – 0,36454	loss – 0,33134
ICDAR 2011 SigVer	acc – 0,635	acc – 0,547	acc – 0,655	train: acc – 0,992 loss – 0,01589	acc – 0,816
	loss – 0,38700	loss – 0,49077	loss – 0,38072	test: acc – 0,993 loss – 0,00757	loss – 0,22358
Own dataset № 1					

от того, является ли она подлинной или поддельной. После этого архитектура немного корректировалась, а ИНС проходила дообучение по описанной выше схеме¹¹.

Все модели обучали в течение 45–75 циклов, в качестве функции потерь использовали CosineEmbeddingLoss (косинусные потери при встраивании). Для исключения эффекта переобучения ИНС перед каждой итерацией незначительно изменяли изображения подписи в виде случайного поворота в пределах 17° в каждую сторону и немного уменьшали их в пределах от 1 (исходный масштаб) до 0,85 (уменьшенный масштаб).

3.3. Анализ полученных результатов

Как можно видеть из результатов, представленных в таблице 2, на всех наборах данных для обучения моделей результаты точности на обучающей и тестовой выборке составляли 74,6–99,7 % (в среднем 92,0 %). Относительно низкую точность в наборе данных GPDS150 (79,5% на обучающей и 74,6% на тестовой выборке) можно объяснить большим количеством авторов, и как следствие, большим разнообразием образцов подписи. Это свидетельствует о необходимости дальнейшего совершенствования архитектуры ИНС для повышения точности модели при ее обучении на тех наборах данных, где имеется большое разнообразие авторов и, как следствие, подписей.

В то же время при тестировании на dataset, на которых модели не проходили обучение, отмечалось существенное снижение точности. Так, если модель была обучена на CEDAR, то показатели точности для других наборов данных, выступающих в этом случае как тестовые, снижались почти на 45%. Это неудивительно, поскольку в этом dataset содержатся образцы подписей, предоставленные лишь 55 авторами, чего явно недостаточно для обобщения. Схожий результат получили Д.В. Бахтеева и Р. Сударикова [35], которые в своей работе проверяли эффективность нейросети SigNet, разработанной Центром компьютерного зрения Автономного университета Барселоны [23], на собранном ими наборе данных подписей, выполненных на кириллице. В ходе экспериментов они обнаружили, что при анализе этого набора данных общая точность модели снизилась на 14%: 94,37% при обучении модели на CEDAR и $80,87 \pm 1.39\%$ на NSP соответственно.

Возвращаясь к результатам данного исследования, отметим, что наибольший средний показатель точности (78,7 %) по всем используемым в работе dataset был получен моделью, обученной на наборе ICDAR 2011 SigVer, где содержатся рукописные подписи 69 авторов, а наименьшим стандартным отклонением в точности (6%) среди всех наборов данных характеризовалась модель, обученная на GPDS150, где общая точность на тестовой выборке составляла только 74,6 % при 150 авторах. Это подтверждает предположение о необходимости совершенствовании архитекту-

¹¹ Исключение составляет только пункт 2 – в нем пара подбиралась как для оригинальных подписей, так и для «подделок», поскольку известно, кем они были выполнены.

Таблица 2. Результаты тестирования искусственной нейронной сети
Table 2. Artificial neural network test results

Тестовые данные \ Обучающие данные	CEDAR	GPDS150	ICDAR 2011 SigVer	Own dataset №1	Own dataset №2	Среднее значение по всем наборам данных	Стандартное отклонение точности	Сравнение точности по отношению к набору данных № 1	Сравнение точности по отношению к набору данных № 2
CEDAR	0,997	0,526	0,558	0,711	0,781	0,714	0,189	0,286	0,216
GPDS150	0,690	0,746	0,621	0,675	0,776	0,701	0,060	0,071	-0,03
ICDAR 2011 SigVer	0,869	0,6	0,945	0,720	0,801	0,787	0,133	0,225	0,144
Own dataset №1	0,635	0,547	0,655	0,993	0,816	0,729	0,176	-----	0,177

ры ИНС на больших выборках, поскольку в таком случае модели будут лучше анализировать аналогичные dataset с меньшим числом авторов.

При анализе сформированного нами dataset моделями, которые были обучены на других, общедоступных данных в виде CEDAR и т. п., точность прогнозов также снижалась, хоть и не слишком существенно, относительно других результатов: максимум – на 28,6 % для набора № 1 и на 21,6 % для набора № 2, если модель обучена на CEDAR, где всего 55 оригинальных образцов подписей; и минимально на 7 % и +3 % соответственно, если ИНС обучалась на GPDS150. Таким образом, потенциально возможно формировать большие наборы данных в виде образцов оригинальных рукописей и их подделок, обучать на них нейронные сети (SNN в частности), после чего использовать их в «реальных» экспертных исследованиях почерка. Это ставит перед экспертным сообществом не только задачу по формированию подобных баз данных для обучения нейросетей, но и по разработке требований и стандартов по точности таких моделей. Однако нельзя исключать и иных подходов к решению задачи идентификации автора рукописи посредством применения ИНС.

Относительно интерпретируемости разработанной нейросети результаты менее однозначны. Так, если модель была обучена на одном из используемых общедоступных наборов данных, то визуализация карт внимания давала только поверхностное представление о выявленных индивидуальных признаках почерка (рис. 5). В большинстве случаев в качестве «важных» областей помечались практически все элементы подписи. Конечно, даже в этих областях присутствуют наиболее «активные» фрагменты, особенно если анализируются подписи из того dataset, на котором, собственно, и обучалась нейронная сеть. Действительно, при сравнении двух оригинальных подписей (*target label 1*) в большинстве случаев имеются общие зоны активации (области красного цвета), а при сравнении оригинальной и поддельных подписей (*target label -1*) наблюдалась ситуация, когда на одном изображении имеется область активации в определенном месте, а на другом изображении в том же месте ее нет, хотя визуально подписи схожи.

Однако при попытке проделать аналогичную процедуру с нашим набором подписей, получили неудовлетворительный

результат, поскольку, как отмечалось ранее, выделялась вся подпись (рис. 6). Это подтверждает предположение о том, что обучение нейросетевых моделей на dataset, в котором данные организованы только для верификации рукописей, негативно влияет на решение идентификационных задач в части интерпретации полученных результатов.

Намного лучше дело обстояло в случае модели, обученной на нашем собственном наборе данных: наблюдалась более точная локализация областей интереса, и нейросеть выявляла такие признаки, как особенности выполнения росчерка (форма движений при выполнении и соединении элементов письменных знаков, точки пересечения элементов росчерка), особенности размещения начальных элементов подписи (форма движений, частично протяженность и, вероятно, направление движений) как в случае образцов из той части данных, на которых ИНС проходила обучение (рис. 7), так и в случае образцов почерка лиц, на которых нейросеть не обучалась (рис. 8). Причем как и при анализе общедоступных изображений рукописи (по типу CEDAR) образцы подписей, визуально схожие, но относящиеся к разным исполнителям, практически не имели общих областей активации.

В то же время во всех приведенных случаях нейросетью были «пропущены» некоторые важные идентификационные признаки в ряде образцов, среди которых можно выделить: число безбуквенных элементов в подписи, а также их размер, форму отдельных участков и ряд других признаков, например, некоторые точки пересечения движений, их начала и окончания. Кроме того, на изображениях местами можно найти области, которые ИНС обозначила как важные, хотя на самом деле эти фрагменты либо являются случайным шумом, либо имеют небольшое идентификационное значение.

По-видимому, система, в частности модули attention, фокусируются главным образом на геометрических особенностях подписи, в то время как такие признаки, как направление движений при выполнении письменных знаков и их элементов, последовательность отдельных элементов и т. п., вероятно, остаются без внимания или не распознаются. Тем не менее, можно говорить о том, что использование ИНС, обученной на образцах конкретных предполагаемых исполнителей, дает более инфор-

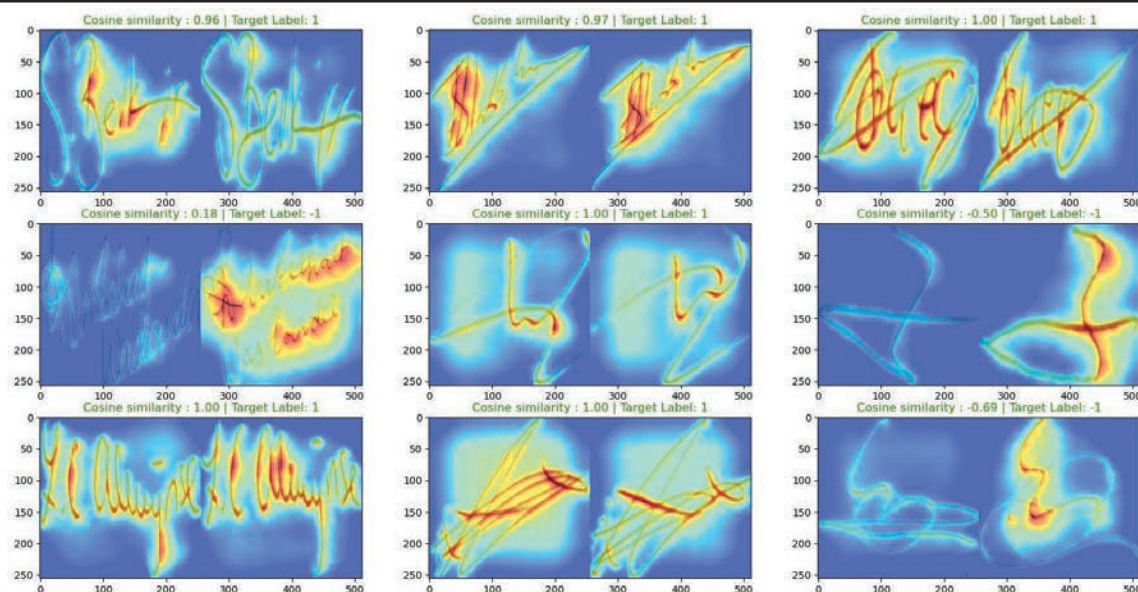


Рис. 5. Визуализация карт внимания для изображений из набора данных ICDAR 2011 SigVer, полученных с использованием модели, обученной на том же наборе данных

Fig. 5. Visualization of attention maps for images from the ICDAR 2011 SigVer dataset obtained using the model trained on the same dataset basis

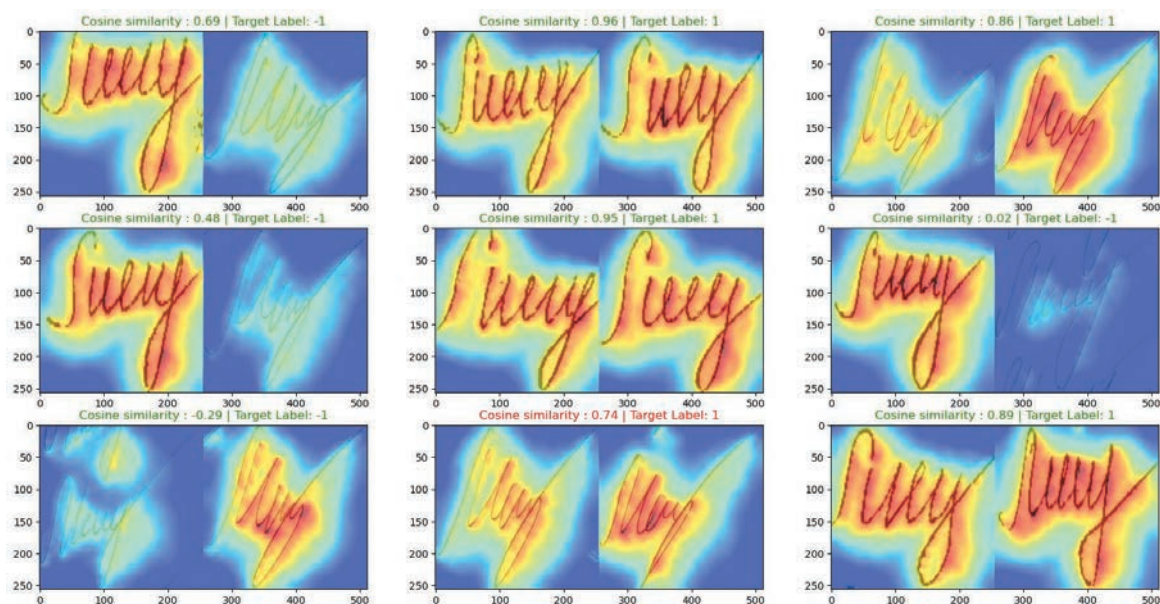


Рис. 6. Визуализация карт внимания для изображений из набора данных, полученных с использованием модели, обученной на данных из ICDAR 2011 SigVer

Fig. 6. Visualization of attention maps for images from the dataset, obtained using a model trained on the data from the ICDAR 2011 SigVer

мативные результаты при интерпретации важных областей рукописей по сравнению с моделями, обученными на общедоступных больших данных. Однако предполагается, что это относится только к системам, основанным в большей степени на использовании сверточных блоков. Вероятнее всего, применение таких подходов, как преобразование рукописного объекта в последовательность штрихов с последующим их ана-

лизом (например, RNN) даст иные результаты.

Закключение

Подводя итог, можно сказать, что, несмотря на все указанные выше проблемы и недостатки примененного в эксперименте подхода конструирования нейросети, по нашему мнению, проведенная работа наглядно показывает, что существующие сегодня ал-

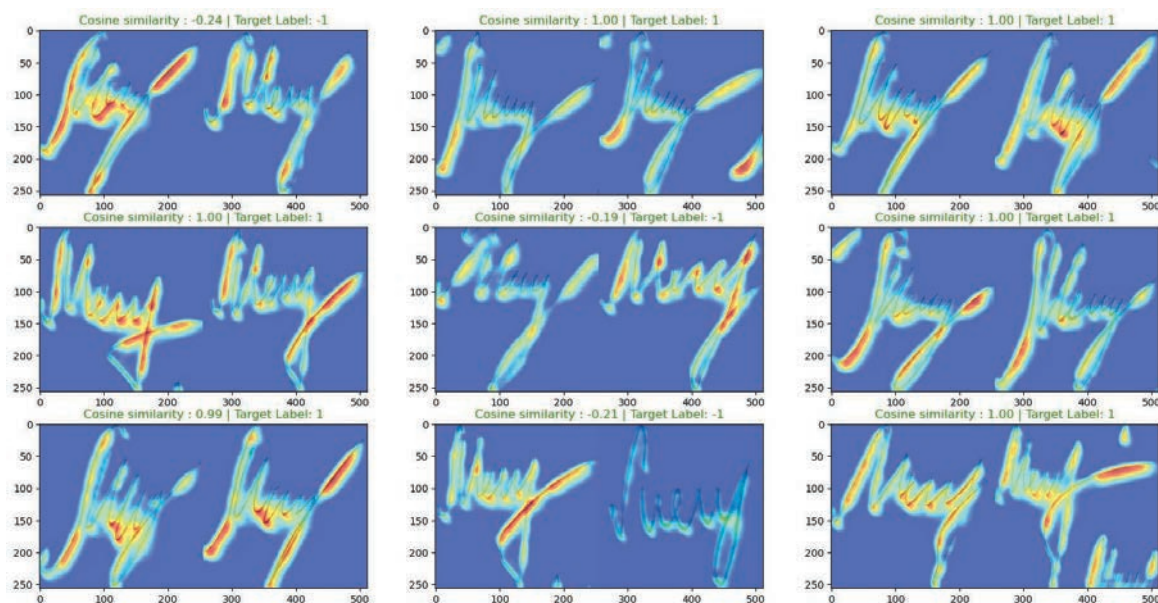


Рис. 7. Визуализация карт внимания для тестовых изображений из подмножества данных, которые были использованы для обучения нейронной сети

Fig. 7. Visualization of attention maps for test images from the subset of data used to train the neural network

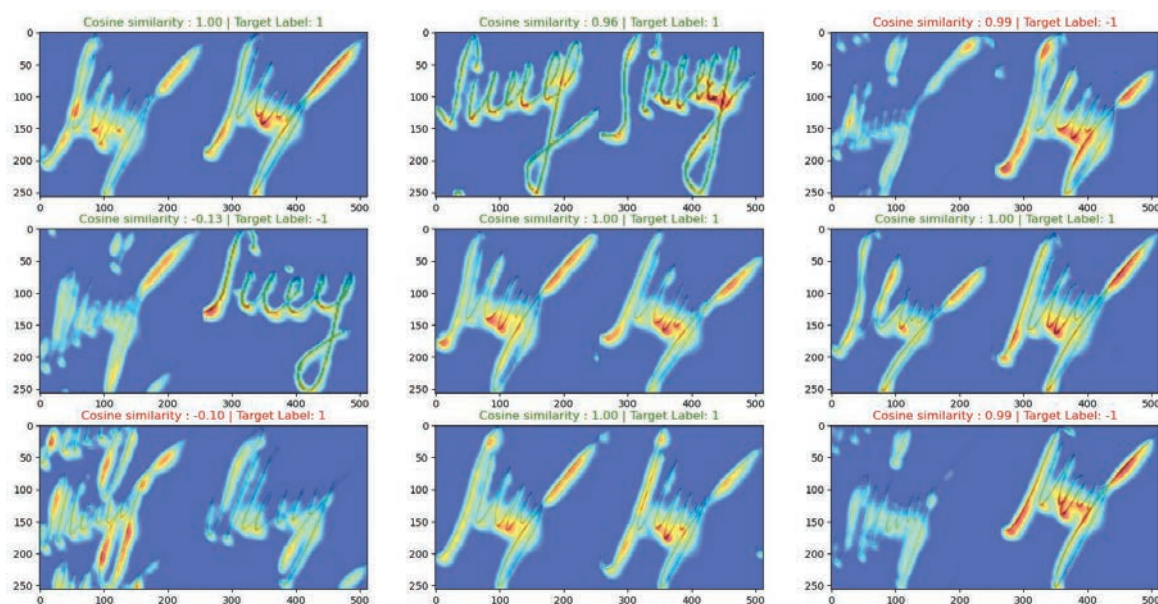


Рис. 8. Визуализация карт внимания для тестовых изображений из подмножества данных, которые не были использованы для обучения нейронной сети

Fig. 8. Visualization of attention maps for test images from the subset of data used to train the neural network

горитмические средства в виде нейронных сетей с механизмом внимания вполне применимы для решения задач судебно-почерковедческой экспертизы, в частности – для идентификации подписи. Вполне вероятно, что при использовании таких архитектур, как Transformer, xLSTM и иных схожих систем можно добиться лучших результатов как в точности прогнозов самой системы, так и в степени интерпретируемости нейросетей.

Что касается описанного подхода интерпретации на основе модуля СВМ – подобные системы могут применяться, к примеру для предварительного анализа почерка с последующим использованием результатов в оперативно-розыскной деятельности, а также в научной работе для поиска новых признаков объектов почерковедческой экспертизы, что особенно актуально в связи с переходом на цифровой документооборот.

Если говорить о решении идентификационных задач при помощи сиамских нейронных сетей, то, на наш взгляд, в настоящее время это наиболее эффективный способ сравнения двух схожих объектов. Как показал эксперимент, а также исследования других авторов, хорошо обученная на большом наборе данных SNN в целом способна справиться со сравнением объектов судебно-почерковедческой экспертизы, даже если она не была обучена на них. Однако, риски возникновения ошибки по-прежнему велики. В связи с этим весьма актуальна разработка методических рекомендаций

по применению сиамской нейронной сети в решении идентификационных задач.

В завершении отметим, что применение в судебно-экспертной деятельности искусственных нейронных сетей и в целом ИИ-алгоритмов, способных за короткий промежуток времени обрабатывать большие массивы данных, становится необходимо. Они позволят значительно упростить и ускорить обработку данных в судебной экспертизе, а также местами автоматизировать наиболее типичные и рутинные этапы экспертного исследования.

СПИСОК ЛИТЕРАТУРЫ

1. Федоренко В.А., Сорокина К.О., Гиверц П.В. Многогрупповая классификация следов бойков с помощью полносвязной нейронной сети // Информационные технологии и вычислительные системы. 2022. № 3. С. 43–57. <https://doi.org/10.14357/20718632220305>
2. Федоренко В.А., Сорокина К.О., Гиверц П.В. Классификация следов бойков по экзemplарам оружия сверточной нейронной сетью // Информационные технологии и вычислительные системы. 2024. № 1. С. 87–96. <https://doi.org/10.14357/20718632240109>
3. Россинская Е.Р. Нейросети в судебной экспертизе: проблемы и перспективы // Вестник Университета имени О.Е. Кутафина (МГЮА). 2024. № 3(115). С. 21–33. <https://doi.org/10.17803/2311-5998.2024.115.3.021-033>
4. Кокин А.В., Денисов Ю.Д. Искусственный интеллект в криминалистике и судебной экспертизе: вопросы правосубъектности и алгоритмической предвзятости // Теория и практика судебной экспертизы. 2023. Т. 18. № 2. С. 30–37. <https://doi.org/10.30764/1819-2785-2023-2-30-37>
5. Андреева О.И., Иванов В.В., Нестеров А.Ю., Трубникова Т.В. Технологии распознавания лиц в уголовном судопроизводстве: проблема оснований правового регулирования использования искусственного интеллекта // Вестник Томского государственного университета. 2019. № 449. С. 201–212. <https://doi.org/10.17223/15617793/449/25>
6. Основы судебной экспертологии: учебно-методическое пособие. М.: ФБУ РФЦСЭ при Минюсте России, 2023. 383 с. <https://doi.org/10.30764/978-5-91133-267-9-2023>
7. Ghorbani A., Abid A., Zou J. Interpretation of Neural Networks Is Fragile // Proceedings of the AAAI Conference on Artificial Intelligence. 2019. Т. 33. No. 1. С. 3681–3688.
8. Das A., Agrawal H., Zitnick L., Parikh D., Batra D. Human Attention in Visual Question Answering: Do Humans and Deep Networks Look at the

REFERENCES

1. Fedorenko V.A., Sorokina K.O., Giverts P.V. Multigroup classification of Firing Pin Marks with the Use of a Fully Connected Neural Network. *Information Technology and Computing Systems*. 2022. No. 3. P. 43–57. (In Russ.). <https://doi.org/10.14357/20718632220305>
2. Fedorenko V.A., Sorokina K.O., Giverts P.V. The Classification of Firing Pin Impressions Using the Convolutional Neural Network (CNN). *Information Technology and Computing Systems*. 2024. No. 1. P. 87–96. (In Russ.). <https://doi.org/10.14357/20718632240109>
3. Rossinskaya E.R. Neural Networks in Forensic Expertology and Expert Practice: Problems and Prospects. *Courier of Kutafin Moscow State Law University (MSAL)*. 2024. Vol. 115. No. 3. P. 21–33. (In Russ.). <https://doi.org/10.17803/2311-5998.2024.115.3.021-033>
4. Kokin A.V., Denisov Yu.D. Artificial Intelligence in Criminalistics and Forensic Examination: Issues of Legal Personality and Algorithmic Bias. *Theory and Practice of Forensic Science*. 2023. Vol. 18. No. 2. P. 30–37. (In Russ.). <https://doi.org/10.30764/1819-2785-2023-2-30-37>
5. Andreeva O.I., Ivanov V.V., Nesterov A.Y., Trubnikova T.V. Facial Recognition Technologies in Criminal Proceedings: Problems of Grounds for the Legal Regulation of Using Artificial Intelligence. *Tomsk State University Journal*. 2019. No. 449. P. 201–212. (In Russ.). <https://doi.org/10.17223/15617793/449/25>
6. *Fundamentals of Forensic Expertology: Educational and Methodological Guide*. Moscow: RFCFS, 2023. 384 p. (In Russ.). <https://doi.org/10.30764/978-5-91133-267-9-2023>
7. Ghorbani A., Abid A., Zou J. Interpretation of Neural Networks Is Fragile. *Proceedings of the AAAI Conference on Artificial Intelligence*. 2019. Vol. 33. No. 1. P. 3681–3688.
8. Das A., Agrawal H., Zitnick L., Parikh D., Batra D. Human Attention in Visual Question Answering: Do Humans and Deep Networks Look at the

- Same Regions? // *Computer Vision and Image Understanding*. 2017. Vol. 163. P. 90–100.
9. Устинов В.В. О возможности объективизации почерковедческого исследования рукописных реквизитов // *Вопросы экспертной практики*. 2019. № S1. С. 663–668.
10. Мещеряков В.А., Бутов В.В. Оценка возможностей почерковедческой экспертизы сквозь призму современных информационных технологий // *Вестник Воронежского института МВД России*. 2017. № 2. С. 40–46.
11. Журавель А.А., Трошко Н.В., Эдзубов Л.Г. Использование алгоритма обобщенного портрета для опознавания образов в судебном почерковедении // *Правовая кибернетика*. 1970. С. 212–227.
12. Woo S., J. Park, Lee J.Y., Kweon I.S. CBAM: Convolutional Block Attention Module // *Proceedings of the European Conference on Computer Vision (ECCV)*. 2018. P. 3–19.
13. Xue G., Liu S., Gong D., Ma Y. ATP-DenseNet: A Hybrid Deep Learning-based Gender Identification of Handwriting // *Neural Computing and Applications*. 2021. Vol. 33. P. 4611–4622.
14. Бобовкин М.В., Диденко О.А., Нестеров А.Е. Компьютерное моделирование раздельного и сравнительного исследования частных признаков почерка на основе программного комплекса «Фрося» // *Судебная экспертиза*. 2021. № 3 (67). С. 62–71. <https://doi.org/10.25724/VAMVD.UXYZ>
15. Zhang X.Y., Xie G.S., Liu C.L., Bengio Y. End-to-End Online Writer Identification with Recurrent Neural Network // *IEEE Transactions on Human-Machine Systems*. 2016. Vol. 47. No. 2. P. 285–292.
16. Kao H.H., Wen C.Y. An Offline Signature Verification and Forgery Detection Method Based on a Single Known Sample and an Explainable Deep Learning Approach // *Applied Sciences*. 2020. Vol. 10. No. 11. 3716. <https://doi.org/10.3390/app10113716>
17. Simonyan K., Vedaldi A., Zisserman A. Deep Inside Convolutional Networks: Visualising Image Classification Models and Saliency Maps // *arXiv*. 2013. <https://doi.org/10.48550/arXiv.1312.6034>
18. Selvaraju R.R., Cogswell M., Das A., Vedantam R., Parikh D., Batra D. Grad-CAM: Visual Explanations from Deep Networks via Gradient-Based Localization // *Proceedings of the IEEE International Conference on Computer Vision*. 2017. P. 618–626.
19. Droby A., Rabaev I., Shapira D.V., Kurar Barakat B., El-Sana J. Digital Hebrew Paleography: Script Types and Modes // *Journal of Imaging*. 2022. Vol. 8. No. 5. P. 143.
20. Marcinowski M. Top Interpretable Neural Network for Handwriting Identification // *Journal of Forensic Sciences*. 2022. Vol. 67. No. 3. P. 1140–1148.
- Same Regions? *Computer Vision and Image Understanding*. 2017. Vol. 163. P. 90–100.
9. Ustinov V.V. On Possibility of Objectivization of Handwriting Expert Examination of Handwritten Requisites. *Issues of Expert Practice*. 2019. No. S1. P. 663–668. (In Russ.).
10. Meshcheryakov V.A., Butov V.V. Evaluation of the Results of the Identical Examination of Script/Writing Due to Modern Information Technologies. *The Bulletin of Voronezh Institute of the Ministry of Internal Affairs of Russia*. 2017. No. 2. P. 40–46. (In Russ.).
11. Zhuravel A.A., Troshko N.V., Edzhubov L.G. Using the Generalized Portrait Algorithm for Pattern Recognition in Forensic Handwriting Examination *Legal Cybernetics*. 1970. P. 212–227. (In Russ.).
12. Woo S., J. Park, Lee J.Y., Kweon I.S. CBAM: Convolutional Block Attention Module. *Proceedings of the European Conference on Computer Vision (ECCV)*. 2018. P. 3–19.
13. Xue G., Liu S., Gong D., Ma Y. ATP-DenseNet: A Hybrid Deep Learning-based Gender Identification of Handwriting. *Neural Computing and Applications*. 2021. Vol. 33. P. 4611–4622.
14. Bobovkin M.V., Didenko O.A., Nesterov A.E. Computer Simulation of Separate and Comparative Study of Private Characteristics Based on the 'Frosya' Software Complex. *Forensic Sciences*. 2021. No. 3 (67). P. 62–71. (In Russ.). <https://doi.org/10.25724/VAMVD.UXYZ>
15. Zhang X.Y., Xie G.S., Liu C.L., Bengio Y. End-to-End Online Writer Identification with Recurrent Neural Network. *IEEE Transactions on Human-Machine Systems*. 2016. Vol. 47. No. 2. P. 285–292.
16. Kao H.H., Wen C.Y. An Offline Signature Verification and Forgery Detection Method Based on a Single Known Sample and an Explainable Deep Learning Approach. *Applied Sciences*. 2020. Vol. 10. No. 11. 3716. <https://doi.org/10.3390/app10113716>
17. Simonyan K., Vedaldi A., Zisserman A. Deep Inside Convolutional Networks: Visualising Image Classification Models and Saliency Maps. *arXiv*. 2013. <https://doi.org/10.48550/arXiv.1312.6034>
18. Selvaraju R.R., Cogswell M., Das A., Vedantam R., Parikh D., Batra D. Grad-CAM: Visual Explanations from Deep Networks via Gradient-Based Localization. *Proceedings of the IEEE International Conference on Computer Vision*. 2017. P. 618–626.
19. Droby A., Rabaev I., Shapira D.V., Kurar Barakat B., El-Sana J. Digital Hebrew Paleography: Script Types and Modes. *Journal of Imaging*. 2022. Vol. 8. No. 5. P. 143.
20. Marcinowski M. Top Interpretable Neural Network for Handwriting Identification. *Journal of Forensic Sciences*. 2022. Vol. 67. No. 3. P. 1140–1148.

21. Harralson H.H., Miller L.S. *Huber and Headrick's Handwriting Identification: Facts and Fundamentals*. CRC press, 2017. 420 p.
22. Bromley J., Guyon I., LeCun Y., Säckinger E., Shah R. Signature Verification using a «Siamese» Time Delay Neural Network // *Advances in Neural Information Processing Systems*. 1993. Vol. 6. P. 669–686.
23. Dey S., Dutta A., Toledo J.I., Ghosh S.K., Llados J., Pal U. SigNet: Convolutional Siamese Network for Writer Independent Offline Signature Verification // *arXiv*. 2017. <https://doi.org/10.48550/arXiv.1707.02131>
24. Bakhteev D.V. Criminalistic Conclusions on Signature Forgery Process During Building an Offline Signature Verification Intellectual System // *Nowa Kodyfikacja Prawa Karnego*. 2021. No. 60. P. 9–15.
25. Chopra P. Making Your Neural Network Say “I Don't Know” – Bayesian NNs using Pyro and PyTorch // *Medium*. <https://towardsdatascience.com/making-your-neural-network-say-i-dont-know-bayesian-nns-using-pyro-and-pytorch-b1c24e6ab8cd>
26. Bahdanau D., Cho K., Bengio Y. Neural Machine Translation by Jointly Learning to Align and Translate // *arXiv*. 2014. <https://doi.org/10.48550/arXiv.1409.0473>
27. Bakhteev D.V., Sudarikov R. NSP Dataset and Offline Signature Verification // *Quality of Information and Communications Technology: 13th International Conference, QUATIC 2020 (Faro, Portugal, September 9–11, 2020)*. Proceedings 13. Springer International Publishing, 2020. P. 41–49.
21. Harralson H.H., Miller L.S. *Huber and Headrick's Handwriting Identification: Facts and Fundamentals*. CRC press, 2017. 420 p.
22. Bromley J., Guyon I., LeCun Y., Säckinger E., Shah R. Signature Verification using a «Siamese» Time Delay Neural Network. *Advances in Neural Information Processing Systems*. 1993. Vol. 6. P. 669–686.
23. Dey S., Dutta A., Toledo J.I., Ghosh S.K., Llados J., Pal U. SigNet: Convolutional Siamese Network for Writer Independent Offline Signature Verification. *arXiv*. 2017. <https://doi.org/10.48550/arXiv.1707.02131>
24. Bakhteev D.V. Criminalistic Conclusions on Signature Forgery Process During Building an Offline Signature Verification Intellectual System. *Nowa Kodyfikacja Prawa Karnego*. 2021. No. 60. P. 9–15.
25. P. Chopra Making Your Neural Network Say “I Don't Know” – Bayesian NNs using Pyro and PyTorch. *Medium*. <https://towardsdatascience.com/making-your-neural-network-say-i-dont-know-bayesian-nns-using-pyro-and-pytorch-b1c24e6ab8cd>
26. Bahdanau D., Cho K., Bengio Y. Neural Machine Translation by Jointly Learning to Align and Translate. *arXiv*. 2014. <https://doi.org/10.48550/arXiv.1409.0473>
27. Bakhteev D.V., Sudarikov R. NSP Dataset and Offline Signature Verification. *Quality of Information and Communications Technology: 13th International Conference, QUATIC 2020 (Faro, Portugal, September 9–11, 2020)*. Proceedings 13. Springer International Publishing, 2020. P. 41–49.

ИНФОРМАЦИЯ ОБ АВТОРЕ

Мищук Всеволод Александрович – аспирант кафедры судебно-экспертной деятельности Юридического института РУДН им. Патриса Лумумбы; e-mail: seva.mi.112@yandex.ru

ABOUT THE AUTHOR

Mishchuk Vsevolod Aleksandrovich – PhD student of the Department of Forensic Science, Institute of Law, RUDN University; e-mail: seva.mi.112@yandex.ru

Статья поступила: 10.03.2025
После доработки: 28.03.2025
Принята к печати: 17.04.2025

Received: March 10, 2025
Revised: March 30, 2025
Accepted: April, 14, 2025