

Научная статья

<https://doi.org/10.29039/2312-7937-2026-2-119-125>

EDN: <https://elibrary.ru/nxtazu>

ПЕТРУХИНА ЕКАТЕРИНА АЛЕКСАНДРОВНА

кандидат юридических наук

Белгородский юридический институт МВД России имени И. Д. Путилина,

(Белгород, Россия)

p.e.a.86@yandex.ru

ORCID: 0009-0000-6475-677X

ЖУРБЕНКО АЛЕКСЕЙ МИХАЙЛОВИЧ

кандидат экономических наук

Всероссийский научно-исследовательский институт МВД России

(Москва, Россия)

zhurbenkoal@yandex.ru,

ORCID: 0000-0001-6022-272X

ПРИМЕНЕНИЕ ИСКУССТВЕННОГО ИНТЕЛЛЕКТА В АНАЛИЗЕ СЛЕДОВ DEERFAKE ТЕХНОЛОГИЙ ПРИ РАССЛЕДОВАНИИ КИБЕРТЕРРОРИЗМА

Аннотация. В статье рассматривается применение искусственного интеллекта для анализа следов deepfake-технологий в расследованиях кибертерроризма. Авторы анализируют алгоритмы искусственного интеллекта на основе нейронных сетей для выявления артефактов в аудио- и видеоматериалах, сгенерированных deepfake. Предложена модель автоматизированного анализа следов, повышающая эффективность досудебного расследования.

Ключевые слова: искусственный интеллект, deepfake, кибертерроризм, анализ следов, нейронные сети, расследование преступлений, цифровая криминалистика

Для цитирования: Петрухина Е. А., Журбенко А. М. Применение искусственного интеллекта в анализе следов deepfake технологий при расследовании кибертерроризма // Вестник ВИПК МВД России. 2026. № 2 (78). С. 119–125; <https://doi.org/10.29039/2312-7937-2026-2-119-125>.

PETRUKHINA EKATERINA A

candidate of law

Belgorod Law Institute of the Ministry of Internal Affairs of the Russian Federation

named after I. D. Putin

(Belgorod, Russia)

ZHURBENKO ALEXEY M.

candidate of economic sciences

National Research Institute of the Ministry of Interior of the Russian Federation

(Moscow, Russia)

THE USE OF ARTIFICIAL INTELLIGENCE IN THE ANALYSIS OF TRACES OF DEEPFAKE TECHNOLOGIES IN THE INVESTIGATION OF CYBERTERRORISM

Abstract. The article examines the use of artificial intelligence to analyze traces of deepfake technology in cyberterrorism investigations. The authors analyze artificial intelligence algorithms based on neural

networks to identify artifacts in audio and video materials generated by deepfake. A model of automated trace analysis has been proposed, which increases the effectiveness of pre-trial investigation.

Keywords: artificial intelligence, deepfake, cyber terrorism, trace analysis, neural networks, crime investigation, digital forensics

For citation: Petrukhina E. A., Zhurbenko A. M. The use of artificial intelligence in the analysis of traces of deepfake technologies in the investigation of cyberterrorism // *Bulletin of the All-Russian Institute of Advanced Training of the Ministry of Internal Affairs of Russia*. 2026. № 2 (78). P. 119–125; <https://doi.org/10.29039/2312-7937-2026-2-119-125>.

В настоящее время киберпространство больше не просто сеть, а полноценное «поле боя», где возникают новые типы угроз. Кибертерроризм среди них необходимо выделить особо.

Один из самых опасных инструментов в руках кибертеррористов – deepfake-технологии. Генеративные нейросети производят синтетический контент, который с пугающей точностью копирует речь, мимику и манеры реальных людей.

Deepfake-технологии создают синтетические медиафайлы. Они имитируют реальное видео, звук или фото с помощью алгоритмов искусственного интеллекта (далее – ИИ).

В основе таких файлов лежат генеративные модели, например GAN или VAE.

В GAN соревнуются две нейросети. Генератор делает подделку, а дискриминатор пытается ее распознать. Так модели повышают свою реалистичность. Обе сети обучаются на больших массивах лиц и голосов. В итоге они умеют заменять лица в видео или синтезировать речь.

Термин «deepfake» образован от deep learning (глубокое обучение) и fake (подделка). Это технология на основе нейросетей, которая выдает очень правдоподобные фальшивые медиа [1, с. 212–217].

«Подделки» эволюционируют до фотореалистичного уровня. VAE, в свою очередь, кодируют данные в латентное пространство, чтобы затем вариативно их восстанавливать. Для видео и звука VAE часто комбинируют с GAN, так алгоритмы анализируют мимику, освещение и произношение. Например, программа DeepFaceLab заменяет лица в роликах, а инструменты voice cloning копируют голос.

Кибертерроризм – незаконные действия в цифровом пространстве. Они угрожают государству, обществу и конкретным людям. Цель – повлиять на политику, экономику или социальные процессы¹.

Deepfake серьезно усиливает возможности кибертеррористов. Они могут создавать пропаганду с аватарами ИИ, фабриковать речи политиков, преувеличивать последствия реальных терактов для громкого резонанса. Такие материалы шокируют общество. Люди теряют доверие к СМИ и властям. Возникает хаос: сбои на выборах, кризисы в дипломатии, страх перед ложными угрозами².

ИИ отлично справляется с большими объемами данных. Он находит скрытые закономерности и аномалии в сетевом трафике, поведении пользователей и системных логах. Это позволяет оперативно обнаружить угрозу – будь то DDoS-атака³ или действия инсайдера.

Так, например, системы машинного обучения выстраивают эталонную модель нормальной работы. Любое отклонение от нее вызывает сигнал тревоги. В итоге время реакции сокращается с часов до секунд. Кроме того, ИИ анализирует историю атак и прогнозирует будущие угрозы. Мониторинг идет в автоматическом режиме, и многие инциденты просто не успевают случиться [2].

В криминалистике ИИ исследует цифровые следы: метаданные файлов, логи доступа, перемещения в сети. Так, эксперты восстанавливают цепочку событий и устанавливают действия злоумышленников.

Методы глубокого обучения распознают дипфейки. Они проверяют артефакты сжатия, неестественную мимику, нестыковки в аудио и видео. Без этого не обойтись, когда нужно опровергнуть фальсификацию по делу о мошенничестве или шантаже.

² Панасенко А. Технологии Deepfake как угроза информационной безопасности // *Anti-Malware.ru*. 08.04.2020. URL: https://www.anti-malware.ru/analytics/Threats_Analysis/Deepfakes-as-a-information-security-threat (дата обращения: 08.04.2026).

³ DDoS-атака (Distributed Denial of Service – «распределенный отказ в обслуживании») – это кибератака, при которой злоумышленники перегружают сервер, сеть или веб-приложение огромным количеством запросов с множества устройств. Цель – сделать систему недоступной для реальных пользователей за счет перегрузки ее ресурсов.

¹ Кибертерроризм // *Anti-Malware.ru*. URL: <https://www.anti-malware.ru/threats/cyberterrorism> (дата обращения: 08.04.2026).

Чтобы отнести инцидент к кибертерроризму, ИИ сопоставляет паттерны атак – их масштаб, цели, возможные мотивы – с базами данных угроз. На выходе получаем классификацию, например по критерию ущерба критической инфраструктуре¹:

- автоматизация работы. ИИ сортирует миллионы событий безопасности. Аналитики освобождаются для стратегических задач;

- система обрабатывает данные в реальном времени и реагирует мгновенно. Ущерб от атаки снижается, а число ложных срабатываний падает – ИИ учится непрерывно;

- масштабируемость. ИИ легко работает с огромными объемами данных. Он встраивается в SIEM и XDR, обеспечивая полное покрытие сетей, облаков и конечных устройств².

Существует несколько методов по распознаванию ИИ:

1. Фотоплетизмография как оптический признак.

Метод удаленной фотоплетизмографии улавливает пульс по едва заметным изменениям цвета кожи. Каждое сокращение сердца меняет кровенаполнение, и алгоритм фиксирует эти микроколебания. Сердце сокращается, концентрация кислорода в коже на мгновение растёт. Красный канал (в RGB или Lab*) тут же меняет свою интенсивность. На реальных видео такие изменения выглядят периодически и когерентно: они синхронно пробегают по лбу, щекам и подбородку.

Синтезированные видео либо вовсе не содержат rPPG-сигнала³, либо выдают хаотичный шум. У этого шума нет характерной частотной структуры в диапазоне 0,6–1,2 Гц (это соответствует пульсу 36–72 удара в минуту).

Алгоритм выявления подделки:

- выделение ROI⁴ – участков лица с густой капиллярной сетью;

- усреднение значения красного канала на этих участках по времени и получение временного ряда;

- применение Фурье- или Вейвлет-анализа⁵, чтобы найти доминирующую частоту или физиологический паттерн;

- сравнение характеристики сигнала: амплитуды, стабильности, согласованности между разными ROI. Deepfake почти всегда выдает рассогласование или полное отсутствие живого ритма.

В контролируемых условиях записи такие методы дают точность 90–95 %. А если добавить к ним другие признаки – точность становится еще выше⁶.

2. Цвет и текстура кожи: гистограммы и пространство Lab*.

Синтезированные лица часто выглядят подозрительно гладкими. Цвет на них переходит слишком плавно, неестественно. Цветовой и текстурный анализ легко это выявляет.

Гистограммы Lab*-каналов живого человека демонстрируют небольшой естественный разброс яркости и оттенков.

У deepfake эти гистограммы словно «прилизаны»: дисперсия ниже, мелкие пиковые отклонения исчезают, цвета становятся подозрительно равномерными.

Анализ текстур обнаруживает, что синтезированные участки кожи менее вариативны и слишком регулярны.

Эти признаки нейросети используют как локальные «мезоскопические» паттерны размером от нескольких до десятков пикселей. Именно на них нацелены детекторы вроде MesoNet (сеть для обнаружения подделок в видео, особенно связанных с манипуляциями лицом)⁷.

3. Микродвижения глаз и мимики.

¹ Цифровые следы, ИИ, биометрия: как исследуют психологию серийных убийц // РБК Тренды. 03.04.2025. URL: <https://trends.rbc.ru/trends/social/67ee283b9a7947314389d27f> (дата обращения: 08.04.2026).

² Искусственный интеллект в сфере кибербезопасности // International academy of cyber security journal, 2025. URL: <https://www.iacj.ru/ru/nauka/article/86419/view> (дата обращения: 08.04.2026).

³ rPPG-сигнал – сигнал, получаемый методом дистанционной фотоплетизмографии (rPPG). Неконтактный способ измерения физиологических параметров, таких как частота сердечных сокращений, частота дыхания, насыщение крови кислородом и артериальное давление, с помощью видеокамер

⁴ ROI (*Region of Interest*) – область интереса, которая в контексте медицинских исследований и диагностики может обозначать конкретные участки тела, подлежащие изучению или анализу.

⁵ Фурье- и Вейвлет-анализ – методы обработки сигналов и данных, которые позволяют преобразовывать сигналы из временной области в частотную или проводить частотно-временной анализ соответственно.

⁶ Новая инфраструктура доверия: технический анализ технологий обнаружения дипфейков и аутентификации цифрового контента // Radensa Newsletter. 30.10.2025. URL: <https://newsletter.radensa.ru/archives/9408> (дата обращения: 08.04.2026).

⁷ Микунов А. В., Елизаров Д. А. Применение модели MesoNet для обнаружения дипфейков // Материалы III конференции (онлайн-ресурс birskin.ru). 18.11.2025. URL: <https://birskin.ru/index.php/2012-02-07-11-31-02/57-iii-1438---mesonet---> (дата обращения: 08.04.2026).

Живое лицо никогда не замирает. Оно постоянно дрожит: тремор глаз, моргание, микросокращения мышц, легкая асимметрия мимики. Большинство deepfake-моделей, особенно в реальном времени, эти детали воспроизводят плохо.

Взрослый человек моргает 15–20 раз в минуту. Длительность закрытия века и симметрия между глазами подчиняются определенной статистике. У deepfake моргание либо отсутствует, либо выглядит механически: слишком регулярно, синхронно, одинаково по длительности, почти никогда не согласуется с другими движениями лица.

Живой глаз слегка и нестабильно дрожит, мы можем отследить это по движению зрачка и контура века. В синтезированных видео таких движений либо нет вовсе, либо траектория выглядит неестественно гладкой и стабильной.

Реальные лица редко двигаются идеально симметрично. Левая и правая половины расходятся – иногда на миллисекунды, иногда по амплитуде. Deepfake-модели, обученные на больших массивах синтезированных кадров, эти различия сглаживают. Они создают «идеально» симметричные движения.

Оптический поток отслеживает движение пикселей между кадрами и выявляет слишком плавные или аномальные траектории.

Facial-landmark-модели фиксируют координаты уголков рта, век, бровей. На выходе – временные ряды движения, которые мы сравниваем с эталонными статистиками живого человека¹.

4. Биометрическая детекция: rPPG, глаза и тепловые карты.

Биометрические методы идут дальше простого rPPG. Они анализируют глаза, тепловое излучение и другие физиологические параметры.

Мы смотрим не только на цвет кожи, но и на изменение диаметра зрачка. Реальный зрачок реагирует на свет, эмоции, когнитивную нагрузку и подчиняется строгим физиологическим законам. В deepfake таких реакций либо нет, либо они заданы жестко и не варьируют.

Кроме того, живое лицо выделяет тепло – от мышц, капилляров, движений. Тепловая карта получается сложной и неоднородной.

У синтетических лиц термальные следы отсутствуют или некорректны, особенно в области глаз, лба и щек.

Натуральная радужка имеет уникальные микроструктуры. Идеально смоделировать их крайне сложно. Детекторы, которые работают с крупным планом, анализируют эти micro-textural-паттерны и отличают живого человека от синтетического двойника.

Такие признаки особенно ценны в системах биометрической аутентификации, где deepfake-спуфинг стал реальной угрозой².

5. 3D-морфинг и проверка глубины.

По видеокадрам алгоритм восстанавливает 3D-меш лица³ или нормализованную карту глубины. В реальном видео глубина и освещение меняются согласованно – вместе с движением источника света и камеры.

Deepfake-методы чаще всего работают в двухмерном пространстве: они накладывают синтезированное лицо на плоский фон. Это порождает рассогласование между картой глубины и реальными тенями, отражениями, движением. Системы, работающие с FaceForensics++⁴, сравнивают локальные 3D-структуры и находят аномалии в форме, контурах и овале лица⁵.

Далее рассмотрим, каким образом ИИ позволяет повысить результативность расследования киберпреступлений, совершенных с использованием deepfake-технологий.

ИИ-системы перерабатывают многотерабайтные массивы цифровых данных – видео, аудио, текст, метаданные. Они выявляют скрытые паттерны, проясняют связи между аккаунтами и находят похожие сценарии распространения экстремистского и террористического контента. Исследо-

² Дипфейки переписывают правила биометрической безопасности // Lazarus Alliance, 2025. URL: <https://lazarusalliance.com/ru/Дипфейки-переписывают-правила-биометрической-безопасности/> (дата обращения: 08.04.2026).

³ 3D-меш лица (Face Mesh) – трехмерная модель человеческого лица, состоящая из сети вершин (точек) и ребер (линий, соединяющих точки), которые определяют форму и контуры лица. Такая модель позволяет точно отслеживать выражения лица, движения и особенности внешности в реальном времени.

⁴ FaceForensics++ – набор данных для обучения методов обнаружения манипулированных изображений лиц. Он был создан в рамках исследования, посвященного анализу реалистичности современных методов манипуляции изображениями и сложности их обнаружения – как автоматически, так и людьми.

⁵ Обзор технологий выявления модифицированного контента класса deepfake // Сбербанк: [сайт]. URL: https://www.sberbank.ru/ru/person/kibrary/experts/obzor_tehnologiy_vyyavleniya_modifitsirovannogo_kontenta_klassa_deepfake (дата обращения: 08.04.2026).

¹ Обзор методов и техник генерации и выявления «дипфейков» // Научно-технический центр ФГУП «ГРЧЦ»: [сайт]. 30.06.2022. URL: https://rdc.grfc.ru/2022/06/deepfakes_generation_method_review/ (дата обращения: 08.04.2026).

ватели подтверждают: связка машинного обучения с цифровой криминалистикой резко ускоряет предварительное расследование и повышает качество экспертизы онлайн-материалов.

Применительно к кибертерроризму ИИ работает в двух режимах: профилактическом и следственном, что позволяет:

- мониторить сеть и находить террористический и экстремистский контент;
- сортировать материалы по степени опасности и принадлежности к группировкам;
- строить версии и предсказывать вероятные цели и методы атак.

Рассмотрим последовательность расследования рассматриваемых нами преступлений:

1. Выявляем deepfake в пропагандистских роликах.

Deepfake-технологии работают на генеративных моделях – GAN, диффузионные сети. Они накладывают чужое лицо, голос или анимацию на исходное видео, и подделка выглядит убедительно. Чтобы ее распознать, нужны специальные инструменты:

- нейросетевые детекторы – они ищут артефакты маскировки: несовпадение движения губ и речи, аномальные тени, странное освещение;
- базы эталонных образцов лиц и голосов (так называемый «биометрический контур»), на которых обучаются модели;
- метаданные – временные метки, геоданные, кодеки, данные об устройстве-источнике. Они помогают отличить смонтированную фальшивку от оригинальной съемки.

Автоматизированные ИИ-сканеры прочесывают крупные видеохостинги и мессенджеры, отсеивая подозрительные ролики. Найденные материалы отправляются на экспертную криминалистическую проверку.

2. Устанавливаем автора и каналы распространения.

На этом этапе ИИ берет на себя три задачи.

- связывает контент с конкретными аккаунтами, доменами и серверами по поведенческим паттернам: частота постов, типы ссылок, сеть ретрансляций;
- обнаруживает скрытые связи между аккаунтами с помощью анализа графов (соцсети, форумы, зашифрованные чаты – через временные и логические пересечения);
- восстанавливает цепочку распространения: от первого источника (пользователь,

канал, сайт) через ретвиты и мессенджеры до конечных ресурсов.

Ключевую роль здесь играют алгоритмы временной и логической корреляции. Они сопоставляют публикацию ролика с другими событиями, например с выступлениями организаций, призывами или инструкциями. Так следователи получают версии о том, какую роль в кампании сыграл тот или иной фактор.

3. Находим связь с экстремистскими и террористическими группировками.

ИИ-модели на основе машинного обучения классифицируют контент по тематическим признакам. Они выделяют:

- призывы к насилию, разжигание ненависти, оправдание терактов;
- ключевые символы, терминологию и нарративы конкретных группировок, например ISIS или «Хаканы национального сопротивления».

Для этого применяются два подхода.

Первый – тематическое моделирование и текстовые классификаторы, обученные на массивах экстремистских материалов.

Второй – эмотивный и тональный анализы¹: оценивают агрессивность, призывность и мобилизующий потенциал текста.

Такие системы автоматически ранжируют контент по степени угрозы и привязывают его к уже известным профилям экстремистских сообществ [3, с. 20–24].

4. Определяем геолокацию и проводим пространственный анализ.

Аналитические алгоритмы ИИ используют три источника:

- геотеги из метаданных изображений и видео;
- контекстные признаки кадра: дорожные знаки, архитектура, ландшафт, транспорт, надписи (их сравнивают с картографическими базами);
- IP-трассировку и данные о серверах-хостингах – так выясняют возможные точки распространения и хранения контента.

На основе этого анализа строятся геопространственные карты распространения. Они показывают, где формируются и откуда ретранслируются пропагандистские ролики. Следствие получает возможность планировать оперативные и международные мероприятия.

¹ Эмотивный и тональный анализы – методы исследования текстов, направленные на выявление эмоциональной окраски, отношения автора к описываемым объектам, событиям или явлениям.

5. Анализируем текст лингвистически.

Лингвистические методы, усиленные ИИ, решают две главные задачи. Первая – установить авторство и стилистические особенности текста: идиолект, характерные лексические и грамматические конструкции. Вторая – определить принадлежность автора к экстремистским сообществам по типичным оборотам, сленгу и метафорам.

Инструментарий включает:

- морфологические и синтаксические анализаторы – они вычленивают структурные паттерны текста;

- количественные модели, которые связывают параметры текста (частота слов, доля эмоциональной лексики, длина предложений) с вероятностью экстремистской или террористической направленности.

Эти алгоритмы не просто автоматически находят подозрительные посты и сообщения. Они еще и формируют предварительные экспертные суждения о социальной и мотивационной вовлеченности автора [4, с. 107–113].

6. Автоматизируем обработку видео.

ИИ-системы здесь дают три ключевых преимущества:

- распознают лица и номера в видеозаписях – даже на низкокачественных и фрагментарных кадрах;

- улучшают качество изображения: суперразрешение, стабилизация, выделение ключевых кадров и сцен;

- создают клип-сводки¹ и временные линии событий на основе видео из разных источников. Это особенно ценно при анализе совместных действий террористических групп.

В криминалистике такой подход эффективно сокращает ручной просмотр – часы, а иногда и сутки видеоматериалов. Эксперты концентрируются только на отобранных и уже проанализированных фрагментах.

7. Строим версии и пути расследования.

ИИ-модели помогают следователям формировать версии развития событий на основе совокупности данных (временные метки, логистика, геолокация, текстовые сообщения). Алгоритмы также оценивают, как контент будет распространяться дальше, какие действия группировки могут предпринять следом.

Необходимо отметить, что ИИ не заменяет следователя. Он работает как аналитический фильтр и эксперт-помощник. Си-

стема предлагает несколько рабочих версий, опираясь на статистические и логические закономерности, а человек уже принимает решение.

8. Выявляем цепочки распространения и точки уязвимости.

Графовые и ML-модели² помогают ИИ находить ключевые узлы в сети распространения. Это аккаунты-дублеры, каналы-ретрансляторы, форумы-агрегаторы. Алгоритмы также распознают типичные паттерны вирусного распространения, например движение из закрытых чатов в публичные каналы или переход от текстовых постов к видеороликам.

Благодаря этому службы безопасности и следственные органы могут:

- оперативно блокировать самые критические точки распространения;

- готовить оперативно-розыскные и технические мероприятия против выявленных узлов сети [5, с. 266–268].

Подводя итоги проведенному исследованию, можно сделать следующие выводы:

1. Цифровизация и развитие ИИ радикально меняют облик киберпреступности. Особенно это заметно в сферах кибертерроризма и дипфейк-угроз.

2. Террористические акты в цифровой среде множатся. Злоумышленники активно используют синтетический контент, чтобы сеять дезинформацию и подрывать доверие к государственным институтам. Киберпространство превратилось в полноценное «поле битвы» – здесь идет информационно-психологическое и политическое противоборство.

3. При расследовании кибертерроризма нейросети берут на себя огромные объемы работы: обрабатывают большие данные, анализируют метаданные, проводят лингвистическую и геопространственную «разведку», выявляют цепочки распространения и строят вероятные сценарии развития событий. Все это освобождает следователя от рутины и ускоряет ход раскрытия и расследования рассматриваемой группы преступлений.

¹ Клип-сводки – тематические видеоматериалы, которые содержат краткие обобщенные сведения или новости по определенной теме.

² Графовые и ML-модели – это два взаимосвязанных подхода для анализа данных, особенно сложных взаимосвязей и структур.

Список источников

1. Севостьянова С. О., Кравцов С. С. Технологии создания deepfake и методы их выявления // Юрист-Правоведъ. 2025. № 2 (113).
2. Газизов А. Р., Хагуш Р. Э. Роль искусственного интеллекта в улучшении обнаружения угроз // Педагогический вестник. 2025. Вып. 36.
3. Квасов В. Ю. Особенности применения ИИ при расследовании преступлений в сфере компьютерных технологий // Парадигма. 2025. № 11.4.
4. Литвинова Т. А., Загоровская О. В. Лингвистические методы выявления в Сети экстремистского контента и лиц, склонных к экстремизму // Современное право. 2016. № 3. URL: <https://www.sovremennoepravo.ru/m/articles/view/Лингвистические-методы-выявления-в-Сети-экстремистского-контента-и-лиц-склонных-к-экстремизму> (дата обращения: 08.04.2026).
5. Вознесенская А. Е. Применение цифровизации в криминалистике: современное состояние и перспективы развития // Молодой ученый. 2024. № 20 (519). URL: <https://moluch.ru/archive/519/114359> (дата обращения: 08.04.2026).

Литература

1. Europol IOCTA cybercrime report reveals growing use of AI models and data theft // Biometric Update. – 2025. – Jun. 16. – URL: <https://www.biometricupdate.com/202506/europol-iocta-cybercrime-report-reveals-growing-use-of-ai-models-and-data-theft> (дата обращения: 08.04.2026).
2. Таршева, М. Н. Актуальные проблемы раскрытия и расследования преступлений террористического характера в цифровую эпоху / М. Н. Таршева, Н. Н. Толкунова // Криминалистика: вчера, сегодня, завтра. – 2025. – № 1.

Информация об авторах:

About the authors:

Петрухина Екатерина Александровна,
доцент кафедры криминалистики;
Журбенко Алексей Михайлович, ведущий
научный сотрудник 3 отдела НИЦ № 3

Petrukhina Ekaterina A., associate professor
of the department of criminalistics;
Zhurbenko Alexey M., leading researcher of
the 3rd department of the Research Center
No. 3

Сведения о вкладе каждого автора

Вклад авторов: все авторы сделали эквивалентный вклад в подготовку публикации.

Авторы заявляют об отсутствии конфликта интересов.

Contribution of the authors: the authors contributed equally to this article.

The authors declare no conflicts of interests.

Статья поступила в редакцию 19.05.2026;

одобрена после рецензирования 09.06.2026; принята к публикации 19.06.2026.

The article was submitted to the editorial office 19.05.2026;

approved after reviewing 09.06.2026; accepted for publication 19.06.2026.