

Кушнир Н. В., Власенко А. В.,
Тютюник Н. П. Смирнов Б. А.

10.14529/secur250402

ПРАКТИЧЕСКИЕ АСПЕКТЫ ПРИМЕНЕНИЯ ИИ-СИСТЕМ ОБНАРУЖЕНИЯ ВТОРЖЕНИЙ

В статье рассматриваются практические аспекты внедрения и эксплуатации ИИ систем обнаружения вторжений в реальных сетях, включая выявление новых угроз, работу в реальном времени и интеграцию с существующей инфраструктурой.

Ключевые слова: ИИ системы обнаружения вторжений, машинное обучение, глубокое обучение, сетевая безопасность, состязательные атаки.

Kushnir N. V. Vlasenko A. V., Tyutyunik N. P. Smirnov B. A.

PRACTICAL ASPECTS OF USING AI-BASED INTRUSION DETECTION SYSTEMS

The article discusses the practical aspects of implementing and operating AI-based intrusion detection systems in real networks, including the identification of new threats, real-time operation, and integration with existing infrastructure.

Keywords: AI-based intrusion detection systems, machine learning, deep learning, network security, and adversarial attacks.

Обнаружение zero-day атак и ранее неизвестных угроз

ИИ-системы IDS способны выявлять атаки, для которых ещё не существуют сигнатуры — zero-day угрозы. Особую роль тут играют методы неконтролируемого и полусконтролируемого обучения, которые фокусируются на аномалиях, отклонениях от нормы и многостадийных атаках [3].

Механизмы выявления аномалий

Применение RNN, LSTM и GAN позволяет фиксировать девиантное поведение без предварительной информации о типах угроз или шаблонов атак. Автокодировщики демонстрируют высокую эффективность в выявлении аномалий: обучаясь только на примерах нормального трафика, они фиксируют значительное увеличение ошибки восстанов-

ления для подозрительных пакетов [4]. Исследования показывают, что автокодировщики способны выявить до 95% новых типов атак, которые не встречались в обучающей выборке.

Практическое значение имеет комбинирование сигнатурных и аномалистических подходов. Система сначала проверяет известные угрозы через быструю сигнатурную базу, а затем направляет остальной трафик через глубокие нейросети для выявления новых паттернов. Такой двухэтапный процесс позволяет значительно снизить количество ложных срабатываний при сохранении высокой чувствительности. В промышленной практике такие системы снижают количество ложных срабатываний на 40-60% в сравнении с использованием только аномалистических подходов.

Анализ поведения пользователей

Профилирование пользовательского поведения (User and Entity Behavior Analytics — UEBA) позволяет выявлять скомпрометированные учётные записи и внутренние угрозы. ИИ-системы анализируют типичные паттерны доступа: рабочее время, типичные источники и целевые системы, объёмы передаваемых данных. При отклонении от этих паттернов система генерирует оповещение. Комбинация UEBA с IDS обеспечивает 360-градусное покрытие угроз как с внешней, так и с внутренней стороны.

Вызовы развёртывания и функционирования в реальном времени

В реальном мире основным ограничением внедрения ИИ является баланс между высокой производительностью и минимальной задержкой обработки информации. Такие системы требуют значительных вычислительных ресурсов: чем сложнее модель, тем выше требования к инфраструктуре. Гигабайты трафика в секунду требуют оптимизированных алгоритмов и архитектуры [5].

Обработка больших объёмов данных

Проблема ложных срабатываний остаётся актуальной и требует регулярного дообучения ИИ на свежих данных, чтобы:

- Уменьшить нагрузку на аналитиков;
- Избежать синдрома «alert fatigue» (усталости от сигналов);
- Адаптировать модель к эволюции сетевого поведения.

В крупных корпоративных сетях ежедневно генерируется от 100 ТБ до 1 ПБ данных сетевого трафика. Традиционные подходы требуют хранения всех этих данных для анализа, что приводит к экспоненциальному росту затрат на инфраструктуру. ИИ-системы позволяют проводить анализ в потоковом режиме (stream processing), выявляя угрозы в режиме реального времени без необходимости сохранения всего трафика.

Оптимизация производительности

Практические решения для оптимизации включают:

1. Кэширование и предвычисления — сохранение результатов часто встречающихся паттернов трафика в памяти высокоскоростных хранилищ (Redis, Memcached). Это может ускорить обработку повторяющихся паттернов на 50-70%.

2. Векторизация и параллелизм — использование GPU (NVIDIA Tesla V100, A100) и TPU (Google TPU) для ускорения матричных операций. Современные ускорители могут обработать трафик со скоростью до 200+ Гбит/с на одном устройстве.

3. Квантизация моделей — снижение точности представления чисел (от float32 к float16 или int8) без значительной потери точности, но с ускорением вычислений [6]. Исследования показывают, что квантизация может ускорить инференс на 3-4 раза с потерей точности менее 1%.

4. Архитектурная оптимизация — использование лёгких моделей для первоначальной фильтрации (например, Random Forest), а сложных моделей (Deep Neural Networks) только для подозрительных пакетов. Это позволяет обработать 80-90% трафика быстро с минимальными вычислительными затратами.

5. Распределённая обработка — использование Spark, Kafka и других систем для параллельной обработки данных на кластерах. Apache Spark может обработать петабайты данных за часы.

Состязательные атаки и устойчивость к обходу

ИИ-модели уязвимы для атак, при которых вредоносный трафик целенаправленно модифицируется для обхода системы обнаружения. Adversarial examples — это минимально изменённые входы, которые вызывают ошибки в классификации [7].

Типы состязательных атак

Исследователи выделяют несколько типов атак на ИИ-системы:

- Evasion attacks — модификация вредоносного трафика на этапе классификации для обхода системы.
- Poisoning attacks — внедрение вредоносных примеров в обучающую выборку для деградации производительности модели.
- Model extraction attacks — попытка скопировать или воспроизвести функциональность чёрного ящика через запросы.

Пример evasion attack: вредонос может добавить в пакет большое количество значений признаков, которые не влияют на его вредоносность, но сдвигают вектор признаков в пространство, где модель его классифицирует как нормальный трафик.

Методы защиты

Для борьбы с состязательными атаками применяются:

- Ансамблевое обучение — использование множества независимых моделей (5-10 различных архитектур) снижает вероятность обхода всех одновременно [8]. Исследования показывают, что ансамбль из 10 моделей на 90% более устойчив к состязательным примерам, чем отдельная модель.

- Состязательное обучение — обучение на специально сгенерированных adversarial примерах, которые пытаются обмануть модель, повышает её устойчивость. Процесс подобен боевой подготовке: система "тренируется" против возможных атак.

- Валидация признаков — проверка логичности значений атрибутов трафика может предотвратить использование нелогичных adversarial примеров. Например, проверка того, что значения портов находятся в диапазоне 0-65535.

- Детектирование аномалий вторичного уровня — дополнительные нейросети, обученные выявлять характеристики adversarial примеров. Это создаёт многоуровневую защиту.

Модели, натренированные с использованием состязательных примеров, демонстрируют повышение устойчивости к обходу на 15-30% в сравнении с базовыми моделями [9].

Интерпретируемость и объяснимость решений

Одна из сложнейших задач внедрения ИИ-IDS — объяснить, почему та или иная сессия классифицирована как угроза. Аналитики безопасности должны понимать, почему система сработала, чтобы принять правильное решение. В продвинутых ИИ-IDS применяются механизмы XAI (Explainable AI) [10]:

Методы объяснимости

- Карты значимости признаков (Feature Importance) — визуализация того, какие атрибуты трафика наиболее повлияли на решение. Например, если система классифицировала пакет как DDoS, она может показать, что решающими факторами были "частота пакетов" (95% вклада) и "размер пакета" (4% вклада).

- LIME (Local Interpretable Model-agnostic Explanations) — построение локальных интерпретируемых моделей вокруг конкретного предсказания [11]. LIME работает путём создания возмущений входных данных и наблюдения, как меняется выход модели.

- SHAP (SHapley Additive exPlanations) — теоретико-игровой подход к определению вклада каждого признака в предсказание модели [12]. SHAP обеспечивает более точные объяснения, чем Feature Importance, благодаря использованию теории кооперативных игр.

- Attention Visualization — для LSTM и Transformer моделей визуализация механизмов внимания показывает, на какие части входа модель обращала внимание при принятии решения. Это может помочь выявить, на какие типы событий реагирует система.

Высокая интерпретируемость критична для принятия решений аналитиками безопасности и для соответствия нормативно-правовым требованиям (GDPR требует объяснения любого решения, влияющего на человека; регулярные аудиты и комплаенс).

Метрики оценки производительности ИИ-IDS

При оценке эффективности ИИ-IDS используются различные метрики:

Основные метрики:

- True Positive Rate (TPR) / Чувствительность (Recall) — доля правильно обнаруженных атак из всех реальных атак. $TPR = TP / (TP + FN)$. Идеальное значение: 100%.

- False Positive Rate (FPR) / Ложных срабатываний — доля неправильно классифицированных нормальных пакетов из всех нормальных пакетов. $FPR = FP / (FP + TN)$. Идеальное значение: близко к 0%.

- Точность (Precision) — доля правильно обнаруженных атак из всех обнаруженных как атаки. $Precision = TP / (TP + FP)$. Высокая точность означает, что система редко выдаёт ложные срабатывания.

- F1-Score — гармоническое среднее между Precision и Recall. $F1 = 2 * (Precision * Recall) / (Precision + Recall)$. Используется для комплексной оценки баланса.

- ROC-AUC (Area Under the Receiver Operating Characteristic Curve) — показатель, объединяющий TPR и FPR. ROC-AUC = 0,5 означает случайное угадывание, ROC-AUC = 1,0 означает идеальную классификацию.

Практические требования

На практике корпорации обычно требуют: $TPR \geq 95\%$ (выявление более 95% атак); $FPR \leq 1\%$ (менее 1% ложных срабатываний); $Detection\ latency \leq 100\text{ мс}$ (обнаружение в течение 100 миллисекунд)

Результаты сравнения

Алгоритм	TPR (%)	FPR (%)	Скорость (пакетов/сек)	Требуемая память
SVM	89	3,2	50	256 МБ
Random Forest	94	2,1	120	512 МБ
KNN	92	2,8	80	1 ГБ
CNN	97	1,5	200	2 ГБ
LSTM	96	1,8	150	3 ГБ
CNN+LSTM	98	0,9	100	4 ГБ
Ансамбль (5 моделей)	99	0,5	60	6 ГБ

Сравнение различных архитектур на практике

Различные архитектуры ИИ показывают разную производительность в зависимости от типа атаки и характера сетевого трафика. Из таблицы 1 видно, что ансамблевые подходы обеспечивают лучший баланс между точностью и ложными срабатываниями, но требуют больше вычислительных ресурсов.

Стоимость внедрения и ROI

При расчёте экономической целесообразности внедрения ИИ-IDS необходимо учитывать различные факторы.

Компоненты стоимости

- 1. Оборудование и инфраструктура — GPU-ускорители (50-100 тыс. руб.), серверы хранения данных (500 тыс.-2 млн. руб.), сетевое оборудование. Типичная стоимость: 2-10 млн. рублей.
- 2. Программное обеспечение и лицензии — лицензии на ИИ-платформы (TensorFlow, PyTorch — открытые), платные платформы для IDS (Fortinet, Palo Alto Networks — от 500 тыс. до 5 млн. руб./год).
- 3. Разработка и внедрение — наём специалистов по ИИ (зарплата 200-400 тыс. руб./месяц), консалтинг, адаптация под специфику сети. Типичная стоимость: 500 тыс.-3 млн. руб.
- 4. Обучение персонала — обучение аналитиков безопасности, операционного персонала. Типичная стоимость: 100-300 тыс. руб.
- 5. Текущее обслуживание и поддержка — ежегодное обновление моделей, техническая поддержка, мониторинг. Типичная стоимость: 10-20% от стоимости внедрения в год.

Расчет ROI

Преимущества ИИ-IDS:

- Предотвращение инцидентов — средняя стоимость инцидента информационной безопасности: 2-20 млн. рублей (в зависимости от типа). Даже предотвращение одного серьёзного инцидента окупает годовые расходы на систему.
 - Снижение затрат на персонал — автоматизация снижает количество необходимых аналитиков с 5-10 до 2-3 человек. Экономия: 1,5-3 млн. руб./год.
 - Сокращение времени отклика — с нескольких часов до минут. Это снижает ущерб от инцидентов на 30-50%.
- Типичный ROI составляет 200-400% в течение первых 3 лет внедрения.

Тестирование и валидация ИИ-IDS

- Перед развёртыванием системы в продакшене необходимо провести комплексное тестирование.
- Фазы тестирования:
- 1. Lab Testing — тестирование на изолированной инфраструктуре с использованием известных датасетов (NSL-KDD, CIC-IDS2017). Целевые метрики: TPR ≥ 95%, FPR ≤ 1%.
 - 2. Staging Testing — развёртывание на копии продакшн-сети. Длительность: 2-4 недели мониторинга без принятия решений (Offline Detection Mode). Целевые метрики: те же, но с реальным трафиком организации.
 - 3. Pilot Deployment — развёртывание на 10-20% сегмента сети в режиме alert-only. Аналитики получают оповещения, но система не блокирует трафик.

4. Full Production Deployment — полное развёртывание с автоматическим блокированием подозрительного трафика.

Типичный период от Lab Testing до Full Deployment: 3-6 месяцев.

Обновление и поддержка системы

ИИ-модели деградируют со временем, так как сетевое поведение эволюционирует, появляются новые типы атак. Это явление называется Data Drift или Concept Drift.

Процесс переобучения:

1. Ежемесячное переобучение — система переучивается на свежих данных, собранных за месяц. Это позволяет учесть изменения в сетевом поведении и новые типы атак.

2. Валидация после переобучения — проверка метрик на холдаут тестовом наборе. Если метрики упали > 5%, переобучение откатывается.

3. A/B тестирование — параллельный запуск старой и новой версий моделей на 5% трафика. Сравнение метрик перед полным развёртыванием.

4. Вероятностное переобучение — использование байесовских методов для оценки неопределённости модели. Если неопределённость растёт, это сигнал к переобучению.

Интеграция с другими системами безопасности

ИИ-IDS не функционирует изолированно, а является частью экосистемы безопасности.

- SIEM (Security Information and Event Management) — корреляция событий от IDS с другими источниками (firewalls, antivirus, logs). Пример: Splunk, IBM QRadar, ArcSight.

- EDR (Endpoint Detection and Response) — интеграция с системами защиты конечных точек. IDS предоставляет контекст о сетевых атаках, EDR анализирует поведение на конечных точках.

- Threat Intelligence Platform — обогащение данных ИИ-IDS информацией о известных угрозах, ботнетах, C2 серверах.

- Firewall и WAF (Web Application Firewall) — автоматическое формирование правил на основе выявленных атак. Пример: если IDS выявил DDoS с определённого IP, firewall блокирует его.

- Ticketing System — автоматическое создание tickets в системе управления инцидентами (Jira, ServiceNow) при выявлении серьёзных атак.

На практике используются API и SIEM connectors для автоматизации этого обмена.

Заключение

ИИ-системы обнаружения вторжений позволяют перейти к проактивной защите сетевых инфраструктур, выявляя как известные, так и ранее неизвестные угрозы. Однако успешное внедрение требует комплексного подхода, включающего: оптимизацию производительности для обработки больших объёмов трафика в реальном времени, защиту от состязательных атак на саму ИИ-систему, обеспечение интерпретируемости и объяснимости решений, регулярное переобучение моделей, интеграцию с существующей инфраструктурой и постоянный мониторинг метрик производительности.

Будущее ИИ-IDS лежит в развитии гибридных моделей, которые объединяют быстрые классические алгоритмы с мощными глубокими нейросетями, в применении механизмов XAI для повышения доверия аналитиков, в разработке методов защиты от состязательных атак и в автоматизации процесса переобучения. Постоянное совершенствование алгоритмов, инвестиции в инфраструктуру и обучение персонала остаются критическими факторами успеха при внедрении таких систем в любой организации.

Литература

1. Ляо, Х. Джей, Линь, К. Х. Р., Линь, И. К., Тунг, К. Й. Система обнаружения вторжений: комплексный обзор // Журнал сетевых и компьютерных приложений. – 2013. – Т. 36, № 1. – С. 16–24. – DOI 10.1016/j.jnca.2012.09.003.
2. Шарафалдин, И., Лашкари, А. Х., Гхорбани, А. А. Создание нового набора данных для обнаружения вторжений и характеристика трафика вторжений // Материалы Международной конференции по информационным системам, безопасности и конфиденциальности (ICISSP). – 2018. – С. 108–116. – DOI 10.5220/0006639001080116.
3. Каусар, А., Жайолд, М., Джиурка, Д. К автоматизированному крупномасштабному обнаружению появляющихся аномалий, независимому от домена // Материалы Международной конференции по искусственному интеллекту и статистике (AISTATS). – 2018. – С. 800–808.
4. Гудфеллоу, И., Бенджио, Й., Курвиль, А. Глубокое обучение / пер. с англ. – М.: ООО "И.Д. Вильямс", 2016. – 800 с.
5. Ким, Дж., Ли, С., Ким, С. Подход глубокого обучения к системе обнаружения вторжений в сети // Материалы Международной конференции по большим данным и смарт-вычислениям (BigComp). – 2016. – С. 313–316. – DOI 10.1109/BIGCOMP.2016.7425916.
6. Чжоу, И., Цзян, Г., Сюй, К. Обнаружение состязательных примеров почти так же сложно, как их классификация // Материалы IEEE Европейского симпозиума по безопасности и конфиденциальности (EuroS&P). – 2019. – С. 250–265. – DOI 10.1109/EuroSP.2019.00023.
7. Гудфеллоу, И. Дж., Шленс, Дж., Сегеди, К. Объяснение и использование состязательных примеров // arXiv препринт. – 2014. – arXiv:1412.6572.
8. Чжоу, З. Х. Методы ансамблевого обучения: основы и алгоритмы / пер. с англ. – М.: CRC Press, 2012. – 290 с.
9. Паперно, Н., МакДэниел, П., Гудфеллоу, И., Джа, С., Челик, З. Б., Свами, А. Практические чёрные атаки на машинное обучение // Материалы 2016 ACM SIGSAC конференции по компьютерной и коммуникационной безопасности. – 2016. – С. 1317–1328. – DOI 10.1145/2976749.2978339.
10. Рибейро, М. Т., Сингх, С., Гвестрин, К. "Почему я должен вам доверять?": Объяснение предсказаний любого классификатора // Материалы 22-й Международной конференции ACM SIGKDD по обнаружению знаний и интеллектуальному анализу данных. – 2016. – С. 1135–1144. – DOI 10.1145/2939672.2939778.
11. Лундберг, С. М., Ли, С. И. Унифицированный подход к интерпретации предсказаний моделей // Advances in Neural Information Processing Systems. – 2017. – Т. 30. – С. 4765–4774.
12. Вавани, А., Шазеер, Н., Пармар, Н., Ушкорейт, Дж., Джонс, Л., Гомес, А. Н., Кайзер, Л., Полозункин, И. Вниманию — это всё, что вам нужно // Advances in Neural Information Processing Systems. – 2017. – Т. 30. – С. 5998–6008.

References

1. Lyao, KH. Dzhey, Lin', K. KH. R., Lin', I. K., Tung, K. Y. Sistema obnaruzeniya vtorzheniy: kompleksnyy obzor // Zhurnal setevykh i komp'yuternykh prilozheniy. – 2013. – Т. 36, № 1. – С. 16–24. – DOI 10.1016/j.jnca.2012.09.003.
2. Sharafaldin, I., Lashkari, A. KH., Gkhorbani, A. A. Sozdaniye novogo nabora dannykh dlya obnaruzeniya vtorzheniy i kharakteristika trafika vtorzheniy // Materialy Mezhdunarodnoy konferentsii po informatsionnym sistemam, bezopasnosti i konfidentsial'nosti (ICISSP). – 2018. – S. 108–116. – DOI 10.5220/0006639001080116.
3. Kausar, A., Zhayold, M., Dzhiurka, D. K avtomatizirovannomu krupnomasshtabnomu obnaruzeniyu poyavlyayushchikhsya anomalii, nezavisimomu ot domena // Materialy Mezhdunarodnoy konferentsii po iskusstvennomu intellektu i statistike (AISTATS). – 2018. – S. 800–808.
4. Gudfellow, I., Bendzhio, Y., Kurvil', A. Glubokoye obucheniye / per. s angl. – M.: ООО "I.D. Vil'yams", 2016. – 800 s.
5. Kim, Dzh., Li, S., Kim, S. Podkhod glubokogo obucheniya k sisteme obnaruzeniya vtorzheniy v seti // Materialy Mezhdunarodnoy konferentsii po bol'shim dannym i smart-vychisleniyam (BigComp). – 2016. – S. 313–316. – DOI 10.1109/BIGCOMP.2016.7425916.
6. Chzhou, I., TSzyan, G., Syuy, K. Obnaruzeniye sostyazatel'nykh primerov pochtii tak zhe slozhno, kak ikh klassifikatsiya // Materialy IEEE Yevropeyskogo simpoziuma po bezopasnosti i konfidentsial'nosti (EuroS&P). – 2019. – S. 250–265. – DOI 10.1109/EuroSP.2019.00023.
7. Gudfellow, I. Dzh., Shlens, Dzh., Segedi, K. Ob'yasneniye i ispol'zovaniye sostyazatel'nykh primerov // arXiv preprint. – 2014. – arXiv:1412.6572.

8. Chzhou, Z. KH. *Metody ansamblevogo obucheniya: osnovy i algoritmy* / per. s angl. – M.: CRC Press, 2012. – 290 s.
9. Paperno, N., MakDeniyel, P., Gudfellou, I., Dzha, S., Chelik, Z. B., Svami, A. *Prakticheskiye chornyye ataki na mashinnoye obucheniye* // *Materialy 2016 ACM SIGSAC konferentsii po komp'yuternoy i kommunikatsionnoy bezopasnosti*. – 2016. – S. 1317–1328. – DOI 10.1145/2976749.2978339.
10. Ribeyro, M. T., Singkh, S., Gvestrin, K. "Pochemu ya dolzhen vam doveryat'?: Ob'yasneniye predskazaniy lyubogo klassifikatora" // *Materialy 22-y Mezhdunarodnoy konferentsii ACM SIGKDD po obnaruzheniyu znaniy i intellektual'nomu analizu dannykh*. – 2016. – S. 1135–1144. – DOI 10.1145/2939672.2939778.
11. Lundberg, S. M., Li, S. I. *Unifitsirovanny podkhod k interpretatsii predskazaniy modeley* // *Advances in Neural Information Processing Systems*. – 2017. – T. 30. – S. 4765–4774.
12. Vasvani, A., Shazeyer, N., Parmar, N., Ushkoreyt, Dzh., Dzhons, L., Gomes, A. N., Kayzer, L., Polozukhin, I. *Vnimaniye — eto vso, chto vam nuzhno* // *Advances in Neural Information Processing Systems*. – 2017. – T. 30. – S. 5998–6008.
-

КУШНИР Надежда Владимировна, старший преподаватель кафедры информационных систем и программирования ФГБОУ ВО «Кубанский государственный технологический университет». 350072, г. Краснодар, ул. Московская, д. 2. E-mail: kushnir.06@mail.ru

ВЛАСЕНКО Александра Владимировна, кандидат технических наук, доцент, Врио заведующего кафедрой информационной безопасности, профессор кафедры ФГКОУ ВО «Краснодарский университет МВД России», 350005, г. Краснодар, ул. Ярославская 128. E-mail: alex_vlasenko@list.ru

ТЮТЮНИК Николай Павлович, студент 3 курса направления подготовки 09.03.04 (программная инженерия) факультета информационных технологий и кибербезопасности (ФИТК) ФГБОУ ВО «Кубанский государственный технологический университет». 350072, г. Краснодар, ул. Московская, д. 2. E-mail: nikolat1236@gmail.com

СМИРНОВ Богдан Александрович, студент 3 курса направления подготовки 09.03.04 (программная инженерия) факультета информационных технологий и кибербезопасности (ФИТК) ФГБОУ ВО «Кубанский государственный технологический университет». 350072, г. Краснодар, ул. Московская, д. 2. E-mail: solo.logger@gmail.com

KUSHNIR Nadezhda Vladimirovna, Senior Lecturer at the Department of Information Systems and Programming of the Federal State Budgetary Educational Institution of Higher Education «Kuban State Technological University». 350072, Krasnodar, Moskovskaya St., 2. E-mail: kushnir.06@mail.ru

VLASENKO Alexandra Vladimirovna, Candidate of Technical Sciences, Associate Professor, Acting Head of the Department of Information Security Professor of the Department, of the Krasnodar University of the Ministry of Internal Affairs of Russia. 350005, Krasnodar, Yaroslavskaya str., 128. E-mail: alex_vlasenko@list.ru

TYUTYUNIK Nikolay Pavlovich, 3rd year student of the 09.03.04 (Software Engineering) Department of Information Technology and Cybersecurity (FITK) of the Federal State Budgetary Educational Institution of Higher Education «Kuban State Technological University». 350072, Krasnodar, Moskovskaya St., 2. E-mail: nikolat1236@gmail.com

SMIRNOV Bogdan Alexandrovich, 3rd year student of the 09.03.04 training course (software Engineering), Faculty of Information Technology and Cybersecurity (FITC), Federal State Budgetary Educational Institution of Higher Education «Kuban State Technological University». 350072, Krasnodar, Moskovskaya St., 2. E-mail: solo.logger@gmail.com