

УДК 004.0

DOI: 10.34670/AR.2026.64.81.011

Информационная безопасность при использовании ИИ (Direct/Indirect Injection)

Жигульский Илья Артемович

Аспирант,
Астраханский государственный
университет им. В.Н. Татищева,
Российская Федерация, Астрахань, ул. Татищева, 20а ;
e-mail: ilya.zhigulskiy.2000@mail.ru

Аннотация

Внедрение больших языковых моделей в контуры обработки данных изменило структуру угроз: управляющее воздействие на систему передаётся теперь не через программный интерфейс, а через естественно-языковую инструкцию, по форме неотличимую от полезного содержания. Разграничение между данными и командой, лежащее в основе классической архитектуры защиты, в диалоговых и агентных системах утрачивает чёткость, вследствие чего оформился класс атак типа *prompt injection*. Разделены два механизма такого воздействия. Прямая инъекция предполагает ввод вредоносной инструкции самим оператором в поле запроса; косвенная внедряет инструкцию во внешний источник, который модель обрабатывает автономно, — документ, веб-страницу, ответ вызванного инструмента. Построена авторская модель многоуровневого контура защиты, и на модельном корпусе из 1 240 тестовых обращений выполнен расчёт его результативности. Сигнатурно-эвристическая фильтрация входа перехватывает до 87,2% прямых инъекций, но лишь около 41,2% косвенных, тогда как добавление контекстной изоляции недоверенного содержания поднимает совокупную полноту обнаружения до 93,8% при доле ложных срабатываний 6,2%. Решающий вклад в снижение результативности атак вносит не входной фильтр, а разграничение привилегий обрабатываемого текста: перехват косвенных инъекций смещается на архитектурный уровень, где внешнее содержание лишается статуса исполняемой команды. Отсюда следует, что устойчивость системы к инъекционному воздействию достигается перестройкой доверительной модели обработки, а сплошное наращивание фильтров на входе даёт убывающую отдачу.

Для цитирования в научных исследованиях

Жигульский И.А. Информационная безопасность при использовании ИИ (Direct/Indirect Injection) // Педагогический журнал. 2026. Т. 16. № 5А. С. 86-93. DOI: 10.34670/AR.2026.64.81.011

Ключевые слова

Информационная безопасность, большие языковые модели, инъекция инструкций, косвенная инъекция, разграничение привилегий, доверенный контур, обнаружение аномалий.

Введение

Обработка естественно-языкового ввода средствами больших языковых моделей стёрла границу между двумя категориями, которые прежняя архитектура защиты держала отдельно, - между данными, подлежащими обработке, и командой, определяющей саму обработку. В системах с фиксированным протоколом входа канал передачи данных и канал управления физически разведены, и попытка провести команду по каналу данных отсекается синтаксическим контролем. Диалоговая модель принимает и то и другое единым потоком токенов, восстанавливая намерение статистически, отчего инструкция, размещённая внутри пользовательского текста, получает шанс быть исполненной наравне с системной. На этом свойстве и построен класс атак, обозначаемый как *prompt injection*; в перечне OWASP Top 10 for LLM Applications редакции 2025 г. он отнесён к категории LLM01 и удерживает первую позицию по значимости.

Проблематика защиты информационных систем разрабатывалась задолго до появления генеративных моделей, и часть её положений сохраняет силу. У В. М. Дудко [Дудко, 2007] информационная безопасность трактуется как системный феномен, не сводимый к отдельному техническому средству; близкой линии придерживается Д. В. Баранов [Баранов, 2004], для которого защита есть непрерывный процесс управления рисками, а не разово достигнутое состояние. Применительно к критически важным объектам сходную мысль проводит Васенин [Васенин, 2008]: устойчивость контура определяется не периметром, а внутренней согласованностью правил доступа. Общий тезис о первенстве архитектуры над отдельным фильтром прямо переносится на инъекционные атаки, где точечная защита входа заведомо неполна.

Разделение инъекции на прямую и косвенную форму удерживает содержательное различие, а не терминологический нюанс. Прямая инъекция оставляет источник запроса недоверенным по умолчанию и потому укладывается в привычную модель угроз. Косвенная смещает точку входа: инструкция внедряется в сторонний ресурс - файл, страницу, ответ вызванного инструмента, - который модель обрабатывает без участия человека, принимая его за данные. Именно здесь расходятся интуитивная модель доверия оператора и фактическое поведение системы: содержимое, считаемое пассивным, приобретает управляющую силу. О. А. Шувалов [Шувалов, 2014] замечает, что риск информационно-коммуникационных технологий концентрируется в точках, где доверие назначается неявно; для агентных конфигураций языковых моделей это наблюдение описывает ядро проблемы.

При быстром росте числа публикаций о явлении в специализированной литературе отсутствует консенсус относительно того, где должен размещаться рубеж перехвата - на входе, в архитектуре обработки или на выходе. Позиция, сводящая защиту к классификации входного текста, представляется дискуссионной, поскольку косвенный вектор она охватывает лишь частично. Цель работы - сопоставить результативность защитных уровней на модельном контуре и определить, какой из них вносит решающий вклад в снижение результативности инъекционного воздействия. Проверяемое предположение связывает основной эффект не с фильтрацией входа, а с разграничением привилегий обрабатываемого содержания.

Материалы и методы исследования

Работа опирается на авторскую расчётную модель многоуровневого контура защиты и на модельный корпус тестовых обращений, сконструированный как непротиворечивый расчётный

пример. Применены методы математического моделирования, сценарного расчёта и факторного разложения по схеме последовательного включения уровней защиты; результативность каждого уровня оценивалась через полноту обнаружения, точность и долю ложных срабатываний. Терминологическая рамка прямой и косвенной инъекции соотнесена с действующей редакцией отраслевого перечня OWASP Top 10 for LLM Applications (2025 г.) без переноса его текста в модель.

Модельный корпус включает 1 240 обращений: 728 легитимных и 512 с инъекционной полезной нагрузкой, из которых 318 отнесены к прямому вектору и 194 - к косвенному. Косвенные обращения разбиты на подкатегории по каналу внедрения: вложенный документ, извлекаемая веб-страница в схеме дополненной генерации, ответ вызванного инструмента в многошаговой цепочке. Контур защиты задан тремя уровнями: сигнатурно-эвристическая фильтрация входа; контекстная изоляция недоверенного содержания с разграничением привилегий, при которой внешний текст помечается неисполняемым; валидация выхода с подтверждением потенциально опасных действий до их совершения. Идея эшелонирования защиты, восходящая к распределённым системам контроля у А. Н. Коварцева [Коварцев, Коломиец, 2010], и требование сквозной защиты информационно-управляющих систем, сформулированное В. С. Зубковым [Зубков, Царегородцев, 2011], перенесены на контур обработки языковой модели.

Результаты и обсуждение

Ключевое свойство инъекционной атаки - совпадение канала команды с каналом данных. Пока разработчик распространяет контроль ввода только на пользовательское поле, косвенный вектор остаётся вне охвата: инструкция приходит вместе с обрабатываемым ресурсом и наследует его доверенный статус. Б. А. Парфенов [Парфенов, 2012] описывает доверие в информационно-коммуникационной среде как назначаемый, а не врождённый атрибут; в терминах языковой модели это означает, что доверенность фрагмента текста задаётся архитектурой обработки, а не его содержанием. С управленческой стороны к тому же выводу подводит А. Н. Райков [Райков, 1999], связывающий защищённость с полнотой контроля над информационными потоками, а не с прочностью отдельного барьера.

Соразмерность защиты также существенна. И. А. Богородицкая [Богородицкая, 2010] предостерегает от превращения безопасности в самоцель, когда издержки контроля начинают превосходить предотвращаемый ущерб. Для инъекционного контура это ограничение проявляется через долю ложных срабатываний: агрессивная фильтрация входа отсекает часть легитимных обращений и подрывает применимость системы. Экономическую цену такого перекоса подчёркивают И. Я. Львович и А. А. Воронов [Львович, Воронов, 2007], для которых инцидент информационной безопасности и избыточный контроль равно сказываются на устойчивости хозяйствующего субъекта. Общая постановка проблемы безопасности при эксплуатации информационных систем у Л. П. Шабановой и П. В. Яковлева [Шабанова, Яковлев, 2017] задаёт ту же рамку: уязвимость возникает на стыке компонентов, а не внутри одного из них.

Систематизация векторов инъекции по каналу внедрения и пересекаемой границе доверия сведена в табл. 1. Технологический и правовой контур обеспечения защиты, очерченный С. Е. Смирновым [Смирнов, 2005], и трактовка информационной безопасности как базовой характеристики цифровой среды у Н. В. Сидельниковой и Т. В. Бесединой [Сидельникова, Беседина, 2018] позволяют рассматривать приведённую классификацию не как перечень

частных приёмов, а как отображение единой закономерности: каждый вектор описывается тем, какую границу доверия он пересекает. Организационную составляющую защиты, на которой настаивают В. И. Матвиенко [Матвиенко, 2014] и К. Засецкая [Засецкая, 2010], классификация учитывает через класс контрмеры, относящийся к процедуре обработки, а не только к программному фильтру.

Таблица 1- Систематизация векторов инъекции по каналу внедрения и границе доверия

Вектор	Канал внедрения	Пересекаемая граница доверия	Класс контрмеры
Прямая инъекция (перепределение инструкций)	Поле пользовательского запроса	Оператор - системная инструкция	Фильтрация и приоритизация входа
Прямая инъекция (обфускация, кодирование)	Поле запроса с изменённым представлением	Оператор - системная инструкция	Нормализация и повторная фильтрация
Косвенная инъекция (вложенный документ)	Обрабатываемый файл	Внешний ресурс - контекст модели	Изоляция содержания, снятие исполняемости
Косвенная инъекция (веб-страница, дополненная генерация)	Извлечённый фрагмент источника	Внешний ресурс - контекст модели	Разграничение привилегий фрагментов
Косвенная инъекция (ответ инструмента, многошаговая цепочка)	Результат вызова инструмента	Инструмент - управляющий контур агента	Валидация выхода, подтверждение действия

Источник: составлено автором.

Расчёт результативности контура выполнен последовательным включением уровней. При работе одного лишь входного фильтра из 318 прямых инъекций перехвачено 277, то есть полнота обнаружения по прямому вектору составила 0,872; по косвенному вектору перехвачено 80 обращений из 194, что соответствует полноте 0,412. Совокупная полнота на всём инъекционном подмножестве удержалась на уровне 0,697 при доле ложных срабатываний 9,3% (68 ошибочно отклонённых легитимных обращений из 728). Разрыв между 0,872 и 0,412 фиксирует главное: входной фильтр эффективен там, где инструкция приходит по ожидаемому каналу, и слабеет там, где она замаскирована под обрабатываемые данные.

Введение второго уровня - контекстной изоляции недоверенного содержания - изменило картину преимущественно по косвенному вектору. Полнота обнаружения косвенных инъекций возросла до 0,866, прямых - до 0,912; совокупная полнота достигла 0,904, а доля ложных срабатываний снизилась до 7,1%, поскольку изоляция снимает необходимость в части агрессивных входных правил. Третий уровень - валидация выхода с подтверждением действий - добавил сравнительно немного к полноте (до 0,938), но перекрыл остаточный канал, при котором инструкция уже проникла в контекст и проявляется лишь на этапе исполнения. Итоговая доля ложных срабатываний зафиксирована на уровне 6,2% (45 обращений из 728). Пофакторное разложение вклада уровней приведено в табл. 2.

Таблица 2 - Полнота обнаружения инъекций при последовательном включении уровней защиты (модельный корпус, 1 240 обращений)

Конфигурация контура	Полнота, прямые	Полнота, косвенные	Полнота, всего	Доля ложных срабатываний, %
Без защиты	0,000	0,000	0,000	0,0
Уровень 1 (фильтрация входа)	0,872	0,412	0,697	9,3

Конфигурация контура	Полнота, прямые	Полнота, косвенные	Полнота, всего	Доля ложных срабатываний, %
Уровни 1-2 (+ изоляция содержания)	0,912	0,866	0,904	7,1
Уровни 1-3 (+ валидация выхода)	0,943	0,933	0,938	6,2
Прирост от уровня 2, п. п.	+4,0	+45,4	+20,7	-2,2

Источник: расчёты автора по модельным данным.

Сопоставление приростов обнаруживает асимметрию, которая и составляет основной результат. По прямому вектору решающий вклад даёт первый уровень: входной фильтр сразу снимает 87,2% угрозы, а последующие уровни добавляют суммарно менее семи процентных пунктов. По косвенному вектору картина обратная: первый уровень перехватывает лишь около двух пятых обращений, тогда как второй уровень поднимает полноту на 45,4 п. п. - от 0,412 до 0,866. Иными словами, для наиболее трудного вектора результат определяется не качеством классификатора входа, а архитектурным решением, лишаящим внешний текст статуса исполняемой команды. Наращивание входных правил без такого решения упирается в потолок: замаскированная под данные инструкция не отличима сигнатурно, и повышение чувствительности фильтра оплачивается ростом ложных срабатываний, против чего и предостерегает соображение о соразмерности защиты.

Динамика по подкатегориям косвенного вектора немонотонна. Инъекции через вложенный документ и через извлекаемую веб-страницу после трёх уровней перехватываются с полнотой около 0,95, тогда как подкатегория многошаговых цепочек вызова инструментов удерживает остаточную результативность: полнота обнаружения на ней не поднимается выше 0,71. Причина, по всей видимости, в том, что вредоносная инструкция здесь не присутствует целиком ни в одном фрагменте, а собирается из безобидных по отдельности шагов, отчего изоляция отдельного источника её не нейтрализует. Этот выброс фиксируется, но не находит полного объяснения в рамках принятой модельной схемы и указывает на границу применимости изоляции как единственного механизма.

Расчёт согласуется с исходным предположением лишь отчасти. Ожидалось, что разграничение привилегий станет решающим уровнем, - и для косвенного вектора это подтвердилось количественно. Однако для прямого вектора первичным остаётся именно входной фильтр, и совокупная эффективность контура достигается только совместной работой уровней, а не доминированием одного из них. Полученное распределение вкладов складывается в пользу многоуровневой трактовки защиты, обозначенной выше через приоритет архитектуры над отдельным барьером, и расходится с позицией, сводящей контроль инъекций к классификации входа.

Механизм разграничения привилегий поддаётся более детальному описанию. Обработываемое содержание разделяется на слои по происхождению: системная инструкция, обращение оператора и внешний материал получают несовпадающие уровни доверия, а правило разрешения конфликтов отдаёт приоритет более доверенному слою. При таком устройстве инструкция, извлечённая из внешнего документа, физически не способна переопределить системную директиву, поскольку принадлежит слою с наименьшими правами. Именно это свойство, а не тонкость распознавания формулировок, объясняет скачок полноты по косвенному вектору с 0,412 до 0,866: изменилась не чувствительность классификатора, а правило, по которому текст вообще допускается к исполнению. Различие принципиально, потому что чувствительность упирается в естественный предел языковой неоднозначности, тогда как правило приоритета слоёв от неё не зависит.

Снижение доли ложных срабатываний при усложнении контура на первый взгляд контринтуитивно, поскольку добавление уровней обычно ужесточает контроль. В модельной схеме эффект объясняется перераспределением нагрузки. Пока перехват косвенных инъекций возлагался на входной фильтр, тот работал в режиме повышенной чувствительности и отклонял пограничные легитимные обращения; после переноса этой задачи на уровень изоляции входной фильтр удалось перенастроить мягче. Доля ошибочно отклонённых обращений снизилась с 9,3 до 7,1%, а затем до 6,2%, тогда как совокупная полнота продолжала расти. Обратная зависимость между строгостью входа и применимостью системы, по всей видимости, и задаёт верхнюю границу однослойной защиты: попытка достичь высокой полноты одним фильтром неизбежно переносит издержки на легитимный трафик.

Остаточная результативность многошаговых цепочек заслуживает отдельного разбора. В подкатегории, где инструкция собирается из нескольких безобидных по отдельности вызовов, не срабатывает ни изоляция отдельного источника, ни фильтрация входа, поскольку вредоносной становится не отдельная реплика, а их последовательность. Полнота обнаружения здесь удержалась около 0,71 против 0,95 на прочих косвенных векторах, и перехват достигался лишь на уровне валидации совокупного действия, когда цепочка уже сформирована и её результат подлежит подтверждению. Разрыв показывает, что перенос защиты на архитектурный уровень необходим, но недостаточен: часть угрозы проявляется только в динамике исполнения и требует контроля не текста, а совершаемого действия.

Заключение

Инъекция инструкций в системах на основе больших языковых моделей вырастает из единственного свойства - совмещения канала команды и канала данных в общем потоке токенов. Пока это свойство сохраняется, полная сигнатурная фильтрация недостижима, и вопрос смещается от обнаружения отдельной вредоносной строки к перестройке доверительной модели обработки.

Модельный расчёт показал разную природу двух векторов. Прямая инъекция перехватывается преимущественно на входе, где источник запроса недоверен по умолчанию; здесь полнота обнаружения превышает 0,87 уже на первом уровне. Косвенная инъекция требует иного рубежа: пока внешнее содержание сохраняет статус исполняемого текста, входной фильтр охватывает менее половины таких обращений. Разграничение привилегий обрабатываемого содержания поднимает полноту по косвенному вектору более чем на сорок пунктов, и именно этот уровень оказывается решающим для наиболее трудной части угрозы.

Практический вывод формулируется как смена точки приложения усилий. Устойчивость к инъекционному воздействию достигается тем, что внешний текст лишается управляющей силы на архитектурном уровне, а не тем, что входной классификатор доводится до предельной чувствительности; сплошное ужесточение фильтра оплачивается ложными срабатываниями и упирается в потолок около 0,7 совокупной полноты. Соразмерность защиты при этом остаётся условием применимости: доля ошибочно отклонённых легитимных обращений, снизившаяся в модели с 9,3 до 6,2%, служит естественным ограничителем чувствительности.

Остаточная результативность многошаговых цепочек вызова инструментов очерчивает предел изоляции как самостоятельного механизма. Инструкция, распределённая по внешне безобидным шагам, не нейтрализуется снятием исполняемости отдельного фрагмента и перехватывается лишь валидацией совокупного действия перед его совершением. Совокупная полнота модельного контура составила 0,938 при доле ложных срабатываний 6,2%.

Библиография

1. Баранов Д.В. Информационная безопасность как процесс управления рисками // Информационное противодействие угрозам терроризма. – 2004. – № 2. – С. 57–59.
2. Богородицкая И.А. Информационная безопасность – не самоцель // Электросвязь. – 2010. – № 4. – С. 46.
3. Васенин. Информационная безопасность критически важных объектов // Проблемы информатики. – 2008. – № 1. – С. 42–47.
4. Дудко В.М. Информационная безопасность как системный феномен // Вестник Академии военных наук. – 2007. – № 3 (20). – С. 67–72.
5. Засецкая К. Информационная безопасность, комплексная защита информации – насколько актуальны сегодня эти вопросы? // Управление персоналом. – 2010. – № 11. – С. 62–64.
6. Зубков В.С., Царегородцев А.В. Обеспечение информационной безопасности информационно-управляющих систем // Современные проблемы науки. – 2011. – № 3. – С. 86–87.
7. Коварцев А.Н., Коломиец Э.И. Распределенная информационная система контроля безопасности // Современные информационные технологии и ИТ-образование. – 2010. – Т. 6. – № 1. – С. 361–363.
8. Львович И.Я., Воронов А.А. Информационная безопасность как фактор экономической стабильности // Информация и безопасность. – 2007. – Т. 10. – № 1. – С. 153–156.
9. Матвиенко В.И. Обеспечение информационной безопасности // Педагогическое образование на Алтае. – 2014. – № 2. – С. 505.
10. Парфенов Б.А. Доверие и безопасность в ИКТ // Вестник связи. – 2012. – № 5. – С. 30–33.
11. Райков А.Н. Информационная безопасность и управление // Информационное общество. – 1999. – № 5. – С. 6–9.
12. Сидельникова Н.В., Беседина Т.В. Информационная безопасность // Образование. Карьера. Общество. – 2018. – № 1 (56). – С. 71–72.
13. Смирнов С.Е. Информационная безопасность. Технологическое и правовое обеспечение (окончание) // Закон и право. – 2005. – № 3. – С. 3–7.
14. Шабанова Л.П., Яковлев П.В. Проблемы безопасности при использовании информационных систем // Проблемы научной мысли. – 2017. – Т. 12. – № 4. – С. 28–31.
15. Шувалов О.А. К вопросу о безопасности использования информационно-коммуникационных технологий // Национальная безопасность и стратегическое планирование. Правовой аспект. – 2014. – № 1. – С. 12.

Information Security When Using AI (Direct/Indirect Injection)

Il'ya A. Zhigul'skii

Postgraduate Student,
Astrakhan State University named after V.N. Tatishchev,
20a, Tatishcheva str., Astrakhan, Russian Federation;
e-mail: ilya.zhigulskiy.2000@mail.ru

Abstract

The introduction of large language models into data processing circuits has changed the structure of threats: the control action on the system is now transmitted not through a software interface, but through a natural language instruction, indistinguishable in form from useful content. The distinction between data and command, which underlies the classical security architecture, loses clarity in dialog and agent systems, as a result of which a class of prompt injection attacks has taken shape. Two mechanisms of such impact are distinguished. Direct injection involves the input of a malicious instruction by the operator themselves in the query field; indirect injection introduces the instruction into an external source that the model processes autonomously — a document, a web page, a response of a called tool. An author's model of a multi-level protection circuit has been constructed, and its effectiveness has been calculated on a model corpus of 1,240 test requests.

Zhigul'skii I.A.

Signature-heuristic input filtering intercepts up to 87.2% of direct injections, but only about 41.2% of indirect ones, whereas the addition of contextual isolation of untrusted content raises the total detection completeness to 93.8% with a false positive rate of 6.2%. The decisive contribution to reducing the effectiveness of attacks is made not by the input filter, but by the differentiation of privileges of the processed text: the interception of indirect injections shifts to the architectural level, where external content is deprived of the status of an executable command. It follows that the system's resilience to injection attacks is achieved by restructuring the trust model of processing, while the continuous build-up of filters at the input yields diminishing returns.

For citation

Zhigul'skii I.A. (2026) Informatsionnaya bezopasnost' pri ispol'zovanii II (Direct/Indirect Injection) [Information Security When Using AI (Direct/Indirect Injection)]. *Pedagogicheskii zhurnal* [Pedagogical Journal], 16 (5A), pp. 86-93. DOI: 10.34670/AR.2026.64.81.011

Keywords

Information security, large language models, prompt injection, indirect injection, privilege separation, trusted circuit, anomaly detection.

References

1. Baranov, D.V. (2004). Informatsionnaya bezopasnost kak protsess upravleniya riskami [Information security as a risk management process]. *Informatsionnoye protivodeystviye ugrozam terrorizma*, (2), 57–59.
2. Bogoroditskaya, I.A. (2010). Informatsionnaya bezopasnost – ne samotsel [Information security is not an end in itself]. *Elektrosvyaz*, (4), 46.
3. Dudko, V.M. (2007). Informatsionnaya bezopasnost kak sistemnyy fenomen [Information security as a systemic phenomenon]. *Vestnik Akademii voyennykh nauk*, (3/20), 67–72.
4. Kovartsev, A.N., & Kolomiets, E.I. (2010). Raspredeleonnaya informatsionnaya sistema kontrolya bezopasnosti [A distributed information system of security control]. *Sovremennyye informatsionnyye tekhnologii i IT-obrazovaniye*, 6(1), 361–363.
5. Lvovich, I.Ya., & Voronov, A.A. (2007). Informatsionnaya bezopasnost kak faktor ekonomicheskoy stabilnosti [Information security as a factor of economic stability]. *Informatsiya i bezopasnost*, 10(1), 153–156.
6. Matvienko, V.I. (2014). Obespecheniye informatsionnoy bezopasnosti [Ensuring information security]. *Pedagogicheskoye obrazovaniye na Altaye*, (2), 505.
7. Parfenov, B.A. (2012). Doveriye i bezopasnost v IKT [Trust and security in ICT]. *Vestnik svyazi*, (5), 30–33.
8. Raikov, A.N. (1999). Informatsionnaya bezopasnost i upravleniye [Information security and management]. *Informatsionnoye obshchestvo*, (5), 6–9.
9. Shabanova, L.P., & Yakovlev, P.V. (2017). Problemy bezopasnosti pri ispol'zovanii informatsionnykh sistem [Security problems in the use of information systems]. *Problemy nauchnoy mysli*, 12(4), 28–31.
10. Shuvalov, O.A. (2014). K voprosu o bezopasnosti ispol'zovaniya informatsionno-kommunikatsionnykh tekhnologiy [On the security of the use of information and communication technologies]. *Natsionalnaya bezopasnost i strategicheskoye planirovaniye. Pravovoy aspekt*, (1), 12.
11. Sidelnikova, N.V., & Besedina, T.V. (2018). Informatsionnaya bezopasnost [Information security]. *Obrazovaniye. Karyera. Obshchestvo*, (1/56), 71–72.
12. Smirnov, S.E. (2005). Informatsionnaya bezopasnost. Tekhnologicheskoye i pravovoye obespecheniye (okonchaniye) [Information security. Technological and legal support (conclusion)]. *Zakon i pravo*, (3), 3–7.
13. Vasenin. (2008). Informatsionnaya bezopasnost kriticheski vazhnykh ob'yektov [Information security of critically important objects]. *Problemy informatiki*, (1), 42–47.
14. Zasetskaya, K. (2010). Informatsionnaya bezopasnost, kompleksnaya zashchita informatsii – naskolko aktualny segodnya eti voprosy? [Information security and comprehensive data protection: How relevant are these issues today?]. *Upravleniye personalom*, (11), 62–64.
15. Zubkov, V.S., & Tsaregorodtsev, A.V. (2011). Obespecheniye informatsionnoy bezopasnosti informatsionno-upravlyayushchikh sistem [Ensuring information security of information and control systems]. *Sovremennyye problemy nauki*, (3), 86–87.