

ЦИФРОВИЗАЦИЯ И УПРАВЛЕНИЕ · DIGITALIZATION AND MANAGEMENT

Вестник МИРБИС. 2026. № 2 (46)'. С. 118–127
Vestnik MIRBIS. 2026; 2 (46)': 118–127

Научная статья
УДК: 005:004.4

DOI: 10.25634/MIRBIS.2026.2.13

Методология управление жизненным циклом самообучающихся цифровых продуктов в условиях динамично меняющихся киберугроз

Светлана Юрьевна Муртузалиева^{1,2}, Артур Валерьевич Бикметов^{1,3}

Аннотация. В статье предлагается методология управления жизненным циклом (ALM) инновационных самообучающихся цифровых продуктов, функционирующих в условиях неопределенности и динамично меняющихся киберугроз. Автор обосновывает переход от классических линейных моделей разработки к кольцевым структурам, учитывающим непрерывное обучение моделей и дрейф данных. Особое внимание уделяется интеграции механизмов объяснимого ИИ (XAI) для повышения прозрачности автономных решений и внедрению концепции «когнитивных двойников» для проактивной защиты инфраструктуры. В работе сформулирована система метрик, объединяющая технические параметры ИИ с операционной эффективностью бизнеса, и определена трансформация роли менеджера продукта в «агентскую эпоху».

Ключевые слова: управление жизненным циклом, самообучающиеся системы, агентский ИИ, кибербезопасность, объяснимый ИИ (XAI), продуктовые метрики, управление данными

Для цитирования: Муртузалиева С. Ю. Методология управление жизненным циклом самообучающихся цифровых продуктов в условиях динамично меняющихся киберугроз / С. Ю. Муртузалиева, А. В. Бикметов. DOI: 10.25634/MIRBIS.2026.2.13 // Вестник МИРБИС. 2026; 2:118–127.

JEL: O32

Original article

Methodology for managing the lifecycle of self-learning digital products in the context of dynamically changing cyber threats

Svetlana Y. Murtuzaliev^{4,5}, Arthur V. Bikmetov^{4,6}

Abstract. The article proposes a lifecycle management methodology (ALM) for innovative self-learning digital products operating under conditions of uncertainty and dynamically changing cyber threats. The author substantiates the transition from classical linear development models to circular structures that account for continuous model training and data drift. Special attention is paid to the integration of explainable AI (XAI) mechanisms to increase the transparency of autonomous decisions and the implementation of the «cognitive twins» concept for proactive infrastructure protection. The paper formulates a system of metrics combining AI technical parameters with business operational efficiency and defines the transformation of the product manager's role in the «agentic era».

Key words: lifecycle management, self-learning systems, agentic AI, cybersecurity, explainable AI (XAI), product metrics, data management

For citation: Murtuzaliev S. Yu. Methodology for managing the lifecycle of self-learning digital products in the context of dynamically changing cyber threats. By S. Yu. Murtuzaliev, A. V. Bikmetov. DOI: 10.25634/MIRBIS.2026.2.13. Vestnik MIRBIS. 2026; 2:118–127 (in Russ.).

JEL: O32

1 Российский университет дружбы народов имени Патриса Лумумбы (РУДН) (Москва, Россия)

2 murtuzaliev-syu@rudn.ru

3 1032241459@rudn.ru

4 Peoples' Friendship University of Russia named after Patrice Lumumba (RUDN University) (Moscow, Russia)

5 murtuzaliev-syu@rudn.ru, <https://orcid.org/0000-0001-9921-8138>

6 1032241459@rudn.ru

Введение

Текущая динамика в сфере информационных технологий свидетельствует о глубинной смене парадигм: индустрия планомерно отходит от разработки систем, функционирующих в рамках жестко детерминированных алгоритмов, в сторону внедрения адаптивных решений с выраженной агентностью. Данная тенденция, ставшая магистральной к середине 2020-х годов, вынуждает экспертное сообщество полностью пересматривать устоявшиеся модели управления жизненным циклом (PLM) и архитектуры информационной безопасности.

Феномен агентского ИИ (Agentic AI) следует рассматривать не как линейное совершенствование алгоритмов машинного обучения, а как качественную трансформацию, при которой системы наделяются функциями автономного целеполагания, операционного планирования и проактивного взаимодействия с внешним инструментарием. Масштабная финансовая экспансия в этот сектор — только за два квартала 2025 года объемом венчурных вливаний в профильные стартапы преодолел отметку в 2,8 млрд долларов — диктует острую необходимость в фундаментальной ревизии существующих управленческих методологий. В настоящее время всё более очевидно, что привычные методы сопровождения интеллектуальных продуктов уже не поспевают за теми технологическими вызовами, которые диктует современная реальность [Апаитов 2023].

Классические модели управления жизненным циклом программного обеспечения — от каскадной Waterfall до итерационных фреймворков семейства Agile — складывались в эпоху, когда системы считались преимущественно детерминированными и их поведение можно было достаточно точно прогнозировать. Однако продукты, способные к самостоятельному обучению, существуют уже в другой реальности: в мире вероятностей, неопределённости и постоянных изменений. Их поведение не остаётся статичным — оно непрерывно трансформируется под влиянием новых данных, поступающих в систему [Лapidус 2021]. Именно поэтому перед теми, кто отвечает за развитие продукта, и перед специалистами по

информационной безопасности встаёт принципиально иная задача. То, что мы выводим в эксплуатацию сегодня, уже через несколько недель реальной работы может демонстрировать совершенно иные поведенческие модели — порой до неузнаваемости [Лapidус 2023].

В ситуации, когда киберугрозы меняются буквально каждый день, а злоумышленники всё активнее применяют агентные системы для автоматизации атак, традиционные средства защиты и контроля перестают справляться. Они просто не обладают достаточной гибкостью, чтобы успевать за такой скоростью перемен.

Главная сложность при формировании новой методологии состоит в том, чтобы найти разумный баланс между автономностью самообучающейся системы и возможностью её осознанного управления. Если раньше классический PLM был сосредоточен главным образом на контроле версий и отслеживании правок в коде, то теперь управление жизненным циклом интеллектуальных продуктов (ALM — AI Lifecycle Management) переносит акцент на работу с данными, постоянный мониторинг «дрейфа» моделей и обеспечение понятности тех решений, которые принимают алгоритмы. К 2025 году стало очевидно: традиционный PLM в своём прежнем виде уже не актуален, он фактически ушёл в прошлое. Ему на смену пришла «эпоха агентов», где умные автономные сущности работают поверх единой цифровой нити (digital thread). Данная концептуальная преемственность выступает в роли интеграционного механизма, обеспечивающего неразрывную конвергенцию антропоцентрических ролей, гетерогенных массивов данных и процессов принятия управленческих решений на всех этапах жизненного цикла высокотехнологичного продукта [Зизаева 2025].

Результаты исследования

Архитектурные детерминанты саморазвивающихся интеллектуальных сред и парадигма нелинейного жизненного цикла

Концептуальная целостность продукта обеспечивается формированием сквозной интегра-

1 © С. Ю. Муртузалиева, А. В. Бикметов, Институт МИРБИС, 2026
Вестник МИРБИС, 2026, № 2(46), с. 118–127.

2 © S. Yu. Murtuzalieva, A. V. Bikmetov, Institute MIRBIS, 2026
Vestnik MIRBIS, 2026, # 2(46), pp. 118–127.

ционной среды, которая выступает связующим звеном между антропоцентрическими ролями, массивами гетерогенных данных и процессами принятия управленческих решений на всех этапах его существования. В современных условиях проектирование самообучающихся цифровых платформ требует внедрения архитектурных решений, способных поддерживать перманентную эволюцию системы без вынужденной стагнации или прерывания операционных процессов. В отличие от традиционного программного обеспечения, функционирование которого ограничено жестко заданными релизами и требует дискретных обновлений, интеллектуальные агенты актуального технологического уклада (2025–2026 гг.) создаются как саморегулирующиеся организмы. Их фундаментальное отличие заключается в способности к инкрементальному обучению на основе поступающих в реальном времени потоков информации, что нивелирует потребность в масштабном ручном вмешательстве со стороны разработчиков [Бусыгин 2025].

Данный архитектурный сдвиг неизбежно влечет за собой отказ от классических линейных моделей в пользу кольцевых или сетевых топологий жизненного цикла. В подобных структурах результаты операционной деятельности немедленно трансформируются в обучающую выборку, создавая замкнутый контур обратной связи, где этап эксплуатации фактически сливается с процессом непрерывного дообучения модели. Методологический базис управления в рамках данной парадигмы выстраивается вокруг системы из семи критических фаз, которые в своей совокупности формируют устойчивый и самоподдерживающийся цикл развития высокотехнологичного продукта. Эти стадии определяют не только технологическую траекторию системы, но и стратегический вектор адаптации продукта к динамически меняющимся требованиям рыночной и кибербезопасной среды. Первая стадия — спецификация сценариев использования и определение бизнес-ценности — требует от команд четкого понимания того, как именно интеллект будет снижать расходы или создавать новые потоки выручки. Вторая стадия — сбор и управление данными — превращается в фундамент всей системы, так как качество прогнозов и отсутствие предвзятости напрямую зависят от чистоты и репрезентативности обучающих выборок. Методо-

логия управления циклом цифровых продуктов представлена в таблице 1.

Таблица 1. Методология управления циклом цифровых продуктов

Этап жизненного цикла	Ключевая задача	Роль ИИ в 2025 году
Идентификация проблемы	определение monetizable-решения	автоматизированный анализ рынка и поиск ниш
Управление данными	очистка, нормализация и аугментация	использование генеративного ИИ для создания синтетических данных
Выбор модели	сравнение фреймворков и архитектур	оценка задержки (latency), объяснимости и стоимости вывода
Разработка логики	построение NLP-пайплайнов и агентов	самодокументируемый код и автоматизированная генерация логики
Тестирование и валидация	контроль галлюцинаций и точности	стресс-тестирование в симуляционных средах (Cognitive Twins)
Развертывание	настройка инфраструктуры мониторинга	автоматическое масштабирование и балансировка нагрузки
Эксплуатация и мониторинг	отслеживание дрейфа модели и обратной связи	непрерывное дообучение (Active Learning) и самокоррекция

Источник: таблица авторов по данным настоящего исследования

Одним из наиболее инновационных подходов в методологии управления является использование «Когнитивных двойников» (Cognitive Twins) — интеллектуальных цифровых реплик экосистемы проекта, которые постоянно учатся на данных для симуляции, прогнозирования и предписания действий в условиях высокой неопределенности, характерной для киберзащиты. Это позволяет продуктовым командам перейти от реактивного исправления ошибок к проактивному моделированию возможных векторов угроз и адаптации продукта до того, как уязвимость будет эксплуатирована [Shneiderman 2020].

Динамический ландшафт киберугроз: адаптация к эпохе генеративного ИИ

Киберугрозы в 2025 году претерпели качественную трансформацию, став более быстрыми,

организованными и автоматизированными. Злоумышленники используют генеративный ИИ для создания полиморфного вредоносного ПО, способного изменять свою сигнатуру при каждом новом заражении, что делает традиционные антивирусные решения бесполезными. В этих условиях методология управления жизненным циклом продукта должна включать «безопасность по умолчанию» (security by design) на уровне архитектуры алгоритмов машинного обучения.

Особое внимание уделяется защите самих ИИ-моделей от специфических атак. К ним относятся инъекции вредоносных промптов (prompt injection), отравление обучающих данных (data poisoning) и состязательные атаки (adversarial manipulation), направленные на то, чтобы заставить модель принять неверное решение путем подачи специально подготовленных входных сигналов. Методология управления должна предусматривать регулярные аудиты на наличие предвзятости и уязвимостей, а также внедрение механизмов состязательного дебайзинга (adversarial debiasing) для обеспечения устойчивости моделей [Гусева 2024]. Основные типы угроз и методология противодействия представлена в таблице 2.

Таблица 2. Основные типы угроз и методология противодействия

Тип угрозы	Механизм воздействия	Методология противодействия
Инъекция промптов	скрытые команды в пользовательском вводе	использование изолированных сред исполнения и фильтрации запросов
Отравление данных	искажение обучающей выборки	робастные пайплайны проверки данных и аудит источников
Полиморфное ВПО	постоянная модификация кода атаки	агентские системы обнаружения на базе поведенческого анализа
Кража модели	извлечение параметров через API-запросы	ограничение частоты запросов и мониторинг аномальных паттернов доступа

Источник: таблица авторов по данным настоящего исследования

В 2025 году критически важным становится внедрение систем непрерывного управления экспозицией угроз (Continuous Threat Exposure

Management, CTEM). Согласно прогнозам экспертов, организации, приоритизирующие безопасность через методологию CTEM, смогут реализовать сокращение числа нарушений безопасности на две трети к 2026 году. Это достигается путем системной оценки достижимости и эксплуатируемости цифровых активов, причем фокус смещается с отдельных компонентов инфраструктуры на бизнес-проекты и векторы угроз в целом [Kotenko 2018].

Методы обеспечения прозрачности и доверия: explainable ai (xai) в кибербезопасности

Одна из самых серьезных проблем при внедрении самообучающихся систем в критически важные сферы, особенно в кибербезопасность, — это так называемый «эффект чёрного ящика». Когда система на основе глубокого обучения вдруг принимает решение изолировать сервер или заблокировать учётную запись высокопоставленного сотрудника, специалисты по безопасности должны чётко понимать, почему именно так произошло. Без этого доверия не будет.

Именно поэтому объяснимый искусственный интеллект (Explainable AI, или XAI) сегодня выходит на первый план в методологии управления жизненным циклом таких систем. Он помогает сделать алгоритмы прозрачными и подотчётными [Çakır 2024].

По данным исследований 2024–2025 годов, примерно две трети экспертов в области информационной безопасности не готовы доверять оповещениям ИИ, если те не сопровождаются внятным объяснением причин [Tian 2025]. Внедрение инструментов объяснимости позволяет не только повысить уровень доверия, но и гораздо быстрее разбираться с ложными срабатываниями.

В современной практике функционирования центров оперативного мониторинга (SOC) наблюдается критический избыток входящих сигналов, что ведет к фактической деградации способности аналитиков своевременно верифицировать угрозы. В этой связи внедрение аналитических инструментов объяснимости, таких как SHAP (Shapley Additive Explanations) и LIME (Local Interpretable Model-agnostic Explanations), становится фундаментальным условием операционной эффективности. Данные методы позволяют декомпозировать логику работы алгоритмов,

выявляя конкретные предикторы, которые определили классификацию события как деструктивного. С точки зрения системного проектирования здесь возникает неизбежная необходимость поиска компромисса между предиктивной мощностью и транспарентностью архитектуры. В настоящее время модели с высокой точностью, как правило, отличаются заметной сложностью интерпретации. При этом резкий переход к существенно более простым архитектурам — например, на основе решающих деревьев — нередко приводит к снижению точности в пределах 15–20 % [Da Silva 2022]. Технологический ландшафт 2025 года явно тяготеет к гибридным решениям. В таких конфигурациях прозрачные модули отвечают за обоснование принимаемых решений, тогда как сложные вычислительные компоненты, работающие по принципу «чёрного ящика», берут на себя основную аналитическую нагрузку. Подобный синтез позволяет сохранять требуемый уровень производительности и одновременно формировать необходимое доверие со стороны корпоративного руководства и государственных регуляторов.

Нормативная база, заданная документами уровня EU AI Act и GDPR, окончательно сделала интерпретируемость обязательным условием эксплуатации интеллектуальных систем. Законодательство прямо требует от организаций обосновывать логику работы автоматизированных решений, особенно когда речь идёт о сферах, затрагивающих права и интересы личности.

В связи с этим методология сопровождения жизненного цикла самообучающихся систем нуждается в серьёзной перестройке. В частности, инструменты объяснимого ИИ (XAI) должны быть встроены непосредственно в процессы проектирования и верификации. На этапе промышленной эксплуатации критически важным становится внедрение удобных интерактивных интерфейсов, которые в реальном времени переводят действия системы на понятный человеку язык. Это даёт операторам возможность оперативно контролировать поведение интеллектуальных агентов и существенно снижает риски, связанные с непрозрачностью алгоритмических выводов [Dasgupta 2020].

Метрики эффективности и математический аппарат управления

Переход к новым подходам в управлении жиз-

ненным циклом самообучающихся систем неизбежно требует иного набора оценочных показателей. Привычные критерии качества, хорошо работавшие для статичного программного обеспечения, в современных условиях оказываются недостаточными, поскольку не учитывают динамическую природу интеллектуальных агентов. Современная парадигма оценки предполагает анализ не только предиктивной точности моделей, но и скорости их адаптации к концептуальному дрейфу данных, а также устойчивости к различным видам деструктивных воздействий, в том числе к состязательным атакам (adversarial attacks).

В рамках такой методологии особое внимание уделяется установлению чётких корреляций между техническими характеристиками модели и прикладными показателями безопасности. Среди ключевых ориентиров здесь выделяются среднее время обнаружения деструктивного события (MTTD) и среднее время восстановления после него (MTTR) [Davenport 2020].

Показатель MTTD отражает оперативность выявления аномалий и рассчитывается как среднее значение временных интервалов между моментом возникновения угрозы и её фиксацией системой:

Показатель MTTD отражает скорость выявления аномалий и рассчитывается как среднее арифметическое временных промежутков от момента возникновения угрозы до её фактической регистрации системой: $MTTD = (\sum(T_{\text{detection}} - T_{\text{occurrence}})/n)$, где $T_{\text{detection}}$ — это момент фиксации инцидента, $T_{\text{occurrence}}$ — точка его реального начала, а n обозначает общее количество зафиксированных случаев.

Параллельно с этим эффективность восстановительных процедур и оперативность реакции оператора оценивается через метрику MTTR, которая записывается следующим образом: $MTTR = (\sum(T_{\text{resolution}} - T_{\text{detection}})/n)$. Здесь $T_{\text{resolution}}$ соответствует полному времени устранения инцидента и возвращения всей инфраструктуры в штатный режим работы.

Интеграция агентного искусственного интеллекта в контур систем оркестрации (SOAR) позволяет достичь существенной оптимизации указанных параметров. Интеллектуальные модули способны в автономном режиме осуществлять высокоуровневую фильтрацию событий, сокра-

щая объем входящего информационного шума на несколько порядков. Это позволяет персоналу центров мониторинга (SOC) фокусировать аналитические ресурсы на верификации действительно критических инцидентов. Валидность использования подобных моделей подтверждается результатами статистического анализа, зафиксированными к 2025 году. В частности, коэффициент корреляции между степенью интеграции ИИ и общей эффективностью процессов составляет $r = 0,816$, в то время как связь с показателем среднего времени наработки на отказ достигает $r = 0,802$, что свидетельствует о высокой степени надежности предлагаемого математического аппарата.

В задачах кибербезопасности регрессионные модели показывают, что комбинированное влияние ИИ и технологий цифровых двойников объясняет до 72 % дисперсии в результатах производительности системы защиты ($R^2 = 0,719$). Основные метрики для управления жизненным циклом самообучающихся продуктов представлены в таблице 3.

Таблица 3. Основные метрики для управление жизненным циклом самообучающихся продуктов

Метрика	Описание	Значение для самообучающихся систем
Precision (Точность)	доля истинно положительных срабатываний среди всех предсказанных	минимизация ложных срабатываний, снижающих «усталость» аналитиков
Recall (Полнота)	доля обнаруженных угроз среди всех реально существующих	гарантия того, что критические атаки не будут пропущены
F0.5 Metric	взвешенное гармоническое среднее с приоритетом точности	баланс между обнаружением и минимизацией операционного шума
Model Drift (Дрейф модели)	изменение статистических свойств данных во времени	индикатор необходимости дообучения или смены архитектуры
Inference Latency	время, затрачиваемое моделью на выдачу предсказания	критический параметр для систем реального времени и защиты сетей

Источник: таблица авторов по данным настоящего исследования

Методология управления жизненным циклом должна включать непрерывный мониторинг этих показателей. При обнаружении дрейфа модели (Model Drift) система должна автоматически инициировать цикл активного обучения (Active Learning), при котором наиболее неопределенные или важные новые данные направляются эксперту-человеку для разметки, после чего модель дообучается [Davenport 2020]. Этот процесс обеспечивает актуальность продукта в условиях постоянно меняющихся тактик и техник злоумышленников (MITRE ATT&CK).

Организационные аспекты и роль менеджера продукта в агентскую эпоху

Переход к самообучающимся продуктам фундаментально меняет роль менеджера продукта (PM). В 2025 году компетенции PM смещаются от простого управления бэклогом функций к стратегическому управлению данными и алгоритмической этике. Менеджер продукта становится «дирижером» агентской экосистемы, который должен не только понимать технические ограничения ИИ, но и уметь интерпретировать сложные аналитические инсайты для принятия бизнес-решений [Li 2018].

Одной из ключевых стратегий управления продуктом становится использование ИИ для автоматизации самой функции менеджмента. Современные платформы управления, такие как ProdPad или Productboard, используют ИИ для скоринга функций на основе обратной связи от клиентов и целей бизнеса (OKRs), что позволяет принимать решения на основе данных, а не интуиции. Это ускоряет время выхода продукта на рынок (Time-to-Market) и повышает точность приоритизации.

В контексте кибербезопасности методология управления жизненным циклом должна интегрироваться с программами изменения поведения и культуры безопасности (SBCEPs). Согласно Gartner, к 2027 году 50 % руководителей информационных служб крупных предприятий внедрят дизайн безопасности, ориентированный на человека, чтобы минимизировать трение между пользователем и инструментами контроля. Это означает, что самообучающиеся продукты должны проектироваться таким образом, чтобы не просто накладывать ограничения, а обучать пользователей принимать правильные решения в области безопасности [Bertiger 2025].

Важной составляющей методологии является выбор модели монетизации и развертывания. В 2025 году доминирует модель SaaS (Software as a Service), которая позволяет провайдеру непрерывно обновлять модели ИИ на своих серверах, обеспечивая всем клиентам мгновенный доступ к новейшим паттернам защиты. Однако для высокочувствительных отраслей сохраняется актуальность гибридных моделей, где обучение происходит в облаке провайдера на деперсонализированных данных, а инференс (вывод) осуществляется локально (on-premise) для обеспечения приватности [Bertiger 2025].

Стратегия «человек в цикле» и этические императивы

Несмотря на то, что агентские системы становятся всё автономнее с каждым годом, в 2025-м управление их жизненным циклом по-прежнему строится вокруг главного правила: человек всегда остаётся в центре процесса (Human-in-the-Loop, или просто HITL). И дело тут не только в технических слабостях ИИ — все мы знаем, как он любит «придумывать» факты, которых на самом деле нет. Не меньшую роль играют этические вопросы и, конечно, юридические риски.

В кибербезопасности это особенно заметно. Там сейчас популярна метафора «костюма Железного человека»: искусственный интеллект словно надевает на аналитика суперспособности — он мгновенно перелопачивает огромные массивы данных, видит то, что человеку не под силу заметить за разумное время. Но в решающий момент, когда, например, нужно срочно отключить целый дата-центр, последнее слово всё равно остаётся за человеком [Рябинин 2025].

С этической точки зрения управление самообучающимися продуктами — это в первую очередь постоянная борьба с предвзятостью алгоритмов. Исследования уже не раз показывали: система обнаружения угроз может вдруг начать считать «вредоносным» трафик только потому, что он идёт из определённого региона или страны. И причина чаще всего кроется в том, на каких данных её когда-то обучали — выборка оказалась неравномерной, и модель «научилась» видеть угрозу там, где её нет.

Поэтому хорошая методология управления обязательно должна включать регулярные аудиты на справедливость (те самые fairness audits) и активно применять техники состязательного

обучения. Только так можно вовремя заметить и нейтрализовать подобные перекосы, пока они не превратились в серьёзную проблему.

Будущее управления жизненным циклом самообучающихся цифровых продуктов лежит в создании адаптивных экосистем, где ИИ-агенты, человеческий интеллект и динамические системы защиты образуют единое целое. К 2026 году ожидается расцвет «коллективной защиты», при которой агенты безопасности из разных организаций смогут обмениваться знаниями о новых паттернах атак в режиме реального времени, создавая глобальный иммунитет против киберпреступности. В этой парадигме управление жизненным циклом превращается из процесса обслуживания продукта в процесс управления непрерывной эволюцией интеллекта.

Интеграция методологии lean six sigma в циклы обучения ИИ

Для обеспечения стабильного качества и минимизации дефектов в самообучающихся продуктах методология управления в 2025 году начинает активно интегрировать принципы Lean Six Sigma. Классический цикл DMAIC (Define, Measure, Analyze, Improve, Control) адаптируется под специфику машинного обучения, превращаясь в инструмент контроля качества данных и точности моделей. В условиях финтеха и кибербезопасности, где цена ошибки крайне высока, использование Six Sigma позволяет отслеживать количество дефектов на миллион возможностей (DPMO) не только в коде, но и в предсказаниях ИИ [IDavenport 2020].

Интеграция ИИ с цифровыми двойниками (Digital Twins) в рамках Six Sigma позволяет организациям достичь значительных улучшений в операционной эффективности. Исследования подтверждают, что такая комбинация обеспечивает статистически значимое снижение уровня дефектов и повышение надежности производства. В управлении продуктом это проявляется в возможности симулировать влияние каждой новой функции на общую стабильность системы еще до ее внедрения в основную ветку разработки [ibid].

Заключение: выводы и стратегические рекомендации

Методология управления жизненным циклом самообучающихся цифровых продуктов в условиях динамично меняющихся киберугроз пред-

ставляет собой многогранный подход, объединяющий передовые достижения в области ИИ, математическую строгость метрик и стратегическую гибкость менеджмента. Основные выводы и рекомендации для экспертов отрасли включают:

1. Имплементация агентской парадигмы в проектировании. Современный подход требует отказа от восприятия цифровых продуктов как статичных наборов функций. Проектирование должно базироваться на создании экосистем автономных агентов, наделенных способностью к самостоятельному операционному планированию и реализации проактивных действий в строгом соответствии с долгосрочными бизнес-целями организации.
2. Инвестиционная приоритизация инфраструктуры данных. Учитывая, что функциональная жизнеспособность интеллектуального продукта на 90% детерминирована качественными характеристиками информационной базы, критически важным становится внедрение жестких протоколов управления данными. Это подразумевает развертывание автоматизированных конвейеров (pipelines) для непрерывной очистки, верификации и обеспечения безопасности входящих потоков.
3. Фундаментальный приоритет интерпретируемости (XAI). Прозрачность алгоритмических выводов выступает базисом доверия в цепочке «система — пользователь» и необходимым условием комплаенса с ак-

туальными регуляторными нормами. Интеграция механизмов объяснимости должна осуществляться не на финальных стадиях, а непосредственно в процессе формирования архитектурного ядра продукта.

4. Переход к непрерывному управлению экспозицией (STEM). В актуальной парадигме информационная безопасность перестает трактоваться как достигнутое статическое состояние. Она трансформируется в перманентный динамический процесс, направленный на идентификацию, оценку и нивелирование потенциальных рисков в режиме реального времени.
5. Гармонизация гибридного взаимодействия «Human-in-the-Loop». Роль эксперта в контуре управления остается преобладающей, особенно в вопросах стратегического целеполагания и этической верификации решений. Искусственный интеллект в данной модели позиционируется как когнитивный ассистент высокого уровня, призванный расширять возможности человека, а не полностью замещать его профессиональные компетенции.

В условиях 2026 года рыночное преимущество в секторе высокотехнологичных продуктов будет принадлежать тем организациям, которые сумеют преобразовать свои системы из пассивных инструментов в активных, эволюционирующих субъектов защиты, способных адаптироваться к вызовам глобального киберпространства быстрее, чем возникают новые векторы атак.

Список источников

- Анаитов, А. М. Управление жизненным циклом цифровых продуктов: методы управления продуктом на разных этапах жизненного цикла — от идеи до снятия с рынка. EDN: OBLVZN // Актуальные вопросы современной экономики = Actual Issues of the Modern Economy. 2023; 8:485–490. eISSN: 2311-4320.
- Бусыгин, Д. Ю. Искусственный интеллект в управлении проектами: возможности и проблемы / Д. Ю. Бусыгин, И. Г. Возмитель. EDN: RIKTJD // Бухгалтерский учет и анализ. 2025; 5:19–23.
- Гусева, М. Н. Перспективы использования искусственного интеллекта в проектном управлении / М. Н. Гусева, И. С. Брикошина, А. И. Глебанов. DOI: 10.34925/EIP.2024.162.1.193. EDN: TZSLCS // Экономика и предпринимательство. 2024; 1:1002–1007. ISSN: 1999-2300.
- Зизаева, А. Искусственный интеллект в системах информационной безопасности / А. Зизаева, З. Ш. Мустафаев. DOI: 10.36871/ek.ur.p.r.2025.09.11.021. EDN: IDVYOD // Экономика и управление: проблемы, решения. 2025; 11(9):186–191. ISSN: 2227-3891 eISSN: 2308-927X.
- Лавидус, Л. В. Стратегии цифровой трансформации бизнеса в условиях нарастающей турбулентности цифровой среды. EDN: VBPBBL // Управление бизнесом в цифровой экономике : Сборник тезисов выступлений Четвертой международной конференции, Санкт-Петербург, 18–19 марта 2021 года. Санкт-Петербург : Санкт-Петербургский государственный университет промышленных технологий и дизайна, 2021. 556 с. С. 20–25. ISBN: 978-5-7937-2101-1.

- Рябинин, А. В. Анализ использования искусственного интеллекта в автоматизированных системах управления / А. В. Рябинин, В. А. Бахметьев. DOI: 10.5281/zenodo.10390725. EDN: REZSLV // Гуманитарный научный вестник. 2023; 11:129–134. eISSN: 2541-7509.
- Shneiderman, B. Human-Centered Artificial Intelligence: Reliable, Safe & Trustworthy. DOI: 10.1080/10447318.2020.1741118 // International Journal of Human-Computer Interaction. 2020; 36(6):1–10.
- Davenport, T. H., Guha, R., Grewal, D. [et al.] How Artificial Intelligence Will Change the Future of Marketing and Product Management. DOI: 10.1007/s11747-019-00696-0 // Journal of the Academy of Marketing Science. 2020; 48(7553):1–19.
- Li, J. H. Cyber security meet artificial intelligence: a survey. DOI: 10.1631/FITEE.1800573 // ENGINEERING Information Technology & Electronic Engineering. 2018; 19(12):1462–1474.
- Dasgupta, D., Akhtar, M. S., Sen, S. Machine Learning in Cybersecurity: A Comprehensive Survey. DOI: 10.1177/1548512920951275 // The Journal of Defense Modeling and Simulation Applications Methodology Technology. 2020; 19(2):2020.
- Bertiger, A., Filar, B., Hastings, J. [et al.] Evaluating LLM Generated Detection Rules in Cybersecurity, 2025. DOI: 10.48550/arXiv.2509.16749. Текст : электронный. URL: <https://arxiv.org/pdf/2509.16749v1> (дата обращения 10.12.2025).
- Çakır, A. M. AI Driven Cybersecurity. DOI: 10.62802/jg7gge06 // Human Computer Interaction. 2024; 8(1):119.
- Da Silva, F. A. Evolving Approaches in Cybersecurity: Metrics and Human Factors. DOI: 10.56238/isevmjv1n2-010 // International Seven Journal of Multidisciplinary. 2022; 1(2). ISSN: 2764-9547.
- Kotenko I. V., Doynikova E., Chechulin A. [et al.] AI- and Metrics-Based Vulnerability-Centric Cyber Security Assessment and Countermeasure Selection: An Artificial Intelligence Approach. DOI: 10.1007/978-3-319-92624-7_5 // Guide to Vulnerability Analysis for Computer Networks and Systems. Springer, 2018. P. 101–130. ISBN: 978-3-319-92623-0.
- Tian, J, Zhu, H. Evaluating the efficacy of AI-driven intrusion detection systems in IoT: a review of performance metrics and cybersecurity threats. DOI 10.7717/peerj-cs.3352 // PeerJ Computer Science. 2025; 11:e3352 DOI 10.7717/peerj-cs.3352.

References

- Apaitov, A. M. Upravleniye zhiznennym tsiklom tsifrovyykh produktov: metody upravleniya produktom na raznykh etapakh zhiznennogo tsikla — ot idei do snyatiya s rynka [Lifecycle Management of Digital Products: Methods of Product Management at Different Stages of the Lifecycle — from Idea to Removal from the Market]. EDN: OBLVZN. *Actual Issues of the Modern Economy*. 2023; 8:485–490. eISSN: 2311-4320 (in Russ.).
- Busygina, D. Yu. Iskusstvennyy intellekt v upravlenii proyektami: vozmozhnosti i problemy [Artificial Intelligence in Project Management: Opportunities and Challenges]. By D. Yu. Busygina, I. G. Vozmitel. EDN: RIKTJD. *Bukhgalterskiy uchët i analiz*. 2025; 5:19–23 (in Russ.).
- Guseva, M. N. Perspektivy ispol'zovaniya iskusstvennogo intellekta v proyektnom upravlenii [Prospects for Using Artificial Intelligence in Project Management]. M. N. Guseva, I. S. Brikoshina, A. I. Glebanov. DOI: 10.34925/EIP.2024.162.1.193. EDN: TZSLCS. *Ekonomika i predprinimatel'stvo*. 2024; 1:1002–1007. ISSN: 1999-2300 (in Russ.).
- Zizayeva, A. Iskusstvennyy intellekt v sistemakh informatsionnoy bezopasnosti [Artificial Intelligence in Information Security Systems]. By A. Zizayeva, Z. Sh. Mustafayev. DOI: 10.36871/ek.up.p.r.2025.09.11.021. EDN: IDVYOD. *Ekonomika i upravleniye: problemy, resheniya*. 2025; 11(9):186–191. ISSN: 2227-3891 eISSN: 2308-927X (in Russ.).
- Lapidus, L. V. Strategii tsifrovoy transformatsii biznesa v usloviyakh narastayushchey turbulentnosti tsifrovoy sredy [Digital Business Transformation Strategies in the Context of Increasing Turbulence in the Digital Environment]. EDN: VBPBBL. *Upravleniye biznesom v tsifrovoy ekonomike* [Business Management in the Digital Economy] : Collection of abstracts from the Fourth International Conference, St. Petersburg, March 18–19, 2021. St. Petersburg : St. Petersburg State University of Industrial Technology and Design Publ., 2021. 556 p. pp. 20–25. ISBN: 978-5-7937-2101-1 (in Russ.).
- Ryabinin, A. V. Analiz ispol'zovaniya iskusstvennogo intellekta v avtomatizirovannykh sistemakh upravleniya [Analysis of the Use of Artificial Intelligence in Automated Control Systems]. By A. V. Ryabinin, V. A. Bakhmetev. DOI: 10.5281/zenodo.10390725. EDN: REZSLV. *Gumanitarnyy nauchnyy vestnik*. 2023; 11:129–

134. eISSN: 2541-7509 (in Russ.).

Shneiderman, B. Human-Centered Artificial Intelligence: Reliable, Safe & Trustworthy. DOI:

10.1080/10447318.2020.1741118. *International Journal of Human-Computer Interaction*. 2020; 36(6):1–10.

Davenport, T. H., Guha, R., Grewal, D. [et al.]. How Artificial Intelligence Will Change the Future of Marketing and Product Management. DOI: 10.1007/s11747-019-00696-0. *Journal of the Academy of Marketing Science*. 2020; 48(7553):1–19.

Li, J. H. Cyber security meet artificial intelligence: a survey. DOI: 10.1631/FITEE.1800573. *ENGINEERING Information Technology & Electronic Engineering*. 2018; 19(12):1462–1474.

Dasgupta, D., Akhtar, M. S., Sen, S. Machine Learning in Cybersecurity: A Comprehensive Survey. DOI: 10.1177/1548512920951275. *The Journal of Defense Modeling and Simulation Applications Methodology Technology*. 2020; 19(2):2020.

Bertiger, A., Filar, B., Hastings, J. [et al.]. *Evaluating LLM Generated Detection Rules in Cybersecurity*, 2025. DOI: 10.48550/arXiv.2509.16749. Online text. URL: <https://arxiv.org/pdf/2509.16749v1> (accessed 12/10/2025).

Çakır, A. M. AI Driven Cybersecurity. DOI: 10.62802/jg7gge06. *Human Computer Interaction*. 2024; 8(1):119.

Da Silva, F. A. Evolving Approaches in Cybersecurity: Metrics and Human Factors. DOI: 10.56238/isevmjv1n2-010. *International Seven Journal of Multidisciplinary*. 2022; 1(2). ISSN: 2764-9547.

Kotenko I. V., Doynikova E., Chechulin A. [et al.] AI- and Metrics-Based Vulnerability-Centric Cyber Security Assessment and Countermeasure Selection: An Artificial Intelligence Approach. DOI: 10.1007/978-3-319-92624-7_5. *Guide to Vulnerability Analysis for Computer Networks and Systems*. Springer, 2018. P. 101–130. ISBN: 978-3-319-92623-0.

Tian, J, Zhu, H. Evaluating the efficacy of AI-driven intrusion detection systems in IoT: a review of performance metrics and cybersecurity threats. DOI 10.7717/peerj-cs.3352. *PeerJ Computer Science*. 2025; 11:e3352 DOI 10.7717/peerj-cs.3352.

Информация об авторах:

Муртузалиева Светлана Юрьевна — кандидат экономических наук, доцент, SPIN-код: 8780-9957, AuthorID: 390020;
Бикметов Артур Валерьевич — магистрант.

Место работы авторов: федеральное государственное автономное образовательное учреждение высшего образования «Российский университет дружбы народов имени Патриса Лумумбы» (РУДН), ул. Миклухо-Маклая, 6, Москва, 117198, Россия. ИНН: 7728073720.

Information about the authors:

Murtuzalieva Svetlana Yu. — Candidate of Economic Sciences, Associate Professor, SPIN-code: 5120-2381, AuthorID: 888069;
Bikmetov Artur V. — master's student.

Affiliation: Peoples' Friendship University of Russia named after Patrice Lumumba (RUDN University), 6 Miklukho-Maklaya St., Moscow, 117198, Russia.

Статья поступила в редакцию 01.04.2026; одобрена после рецензирования 20.04.2026; принята к публикации 26.06.2026.
The article was submitted 04/01/2026; approved after reviewing 04/20/2026; accepted for publication 06/26/2026.