

Мазеин Кирилл Сергеевич

*Студент ФГБОУ ВО «Санкт-Петербургский государственный
технологический институт (технический университет)» Санкт-
Петербург Российская Федерация*

ПРИМЕНЕНИЕ ИСКУССТВЕННОГО ИНТЕЛЛЕКТА В СИСТЕМНОМ АНАЛИЗЕ И ПОДДЕРЖКЕ ПРИНЯТИЯ УПРАВЛЕНЧЕСКИХ РЕШЕНИЙ

Резюме: В работе предложен концептуально-математический подход к включению методов искусственного интеллекта в контур системного анализа и поддержки принятия управленческих решений. Цель исследования состоит в формировании архитектуры интеллектуальной системы поддержки принятия решений, в которой алгоритмические прогнозы отделены от многокритериального выбора и экспертной ответственности лица, принимающего решение. Методическая основа включает системное представление объекта управления, декомпозицию функций интеллектуальной поддержки, нормирование частных критериев, аддитивную свертку и анализ устойчивости рейтинга альтернатив при изменении управленческих приоритетов. Предложена шестислойная архитектура, объединяющая подготовку данных, системное моделирование, прогнозирование, многокритериальную оценку, экспертный контроль и обратную связь. На демонстрационном примере выбора варианта модернизации информационно-управляющей подсистемы показано, что гибридная архитектура, сочетающая машинное обучение с правилами и экспертной проверкой, имеет наибольшую интегральную оценку при сбалансированном наборе критериев. Анализ чувствительности показывает изменение лидирующей альтернативы при существенном повышении веса риска, что подтверждает необходимость явного учета целей и ограничений системы. Полученные результаты могут использоваться при проектировании интеллектуальных систем поддержки решений в организационно-технических системах, где критичны объяснимость,

контролируемость и ответственность за итоговое управленческое воздействие.

Ключевые слова: системный анализ; искусственный интеллект; системы поддержки принятия решений; многокритериальный выбор; машинное обучение; управление; объяснимый искусственный интеллект; человеко-машинное взаимодействие.

Abstract: This paper proposes a conceptual and mathematical approach to incorporating artificial intelligence methods into systems analysis and management decision support. The study aims to develop an architecture of an intelligent decision support system in which algorithmic forecasts are separated from multicriteria choice and from the final responsibility of the human decision maker. The methodology combines a system representation of the controlled object, functional decomposition of intelligent support, normalization of partial criteria, an additive aggregation model, and sensitivity analysis of alternative rankings under changing managerial priorities. A six-layer architecture is proposed that integrates data preparation, system modeling, forecasting, multicriteria assessment, expert control, and feedback. A demonstrative case of selecting an information-management subsystem modernization option shows that a hybrid architecture combining machine learning, rule-based logic, and expert verification achieves the highest integral score under a balanced set of criteria. Sensitivity analysis shows that the leading alternative changes when the risk criterion receives a substantially higher weight, confirming the need to represent system goals and constraints explicitly. The results can be applied to the design of intelligent decision support systems for organizational and technical systems in which explainability, controllability, and accountability for the final management action are essential.

Keywords: systems analysis; artificial intelligence; decision support systems; multicriteria decision-making; machine learning; management; explainable artificial intelligence; human-AI interaction.

Введение

Управление сложными организационно-техническими системами все в большей степени осуществляется в условиях высокой динамики внешней среды, роста числа взаимосвязанных параметров и ограниченного времени на подготовку решения. Для лица, принимающего решение, проблема заключается не столько в дефиците информации, сколько в необходимости выделить из неоднородных потоков данных значимые изменения, определить причинно-следственные связи и оценить последствия возможных управляющих воздействий. В работах по обеспечению ситуационной осведомленности показано, что методы искусственного интеллекта могут сокращать время обработки больших массивов информации и поддерживать идентификацию значимых событий в управленческом контуре [1].

Системы поддержки принятия решений при этом перестают быть исключительно расчетными инструментами. Они становятся человеко-машинными системами, в которых алгоритмические процедуры взаимодействуют с экспертными знаниями, организационными регламентами и ответственностью пользователя. Современные исследования паттернов такого взаимодействия подчеркивают, что качество решения определяется не только точностью вычислительной модели, но и тем, как распределены функции между человеком и машиной, какие данные доступны пользователю и каким образом представлено обоснование рекомендации [2].

В прикладных задачах интеллектуальная поддержка может строиться на различных представлениях знаний. Например, для диагностических систем используются таблицы решений, обеспечивающие явную фиксацию условий, правил и вариантов действий [3]. В более сложных сценариях необходимы онтологические и концептуальные модели, позволяющие связать элементы проблемной ситуации, роли участников и процедуры

совместной деятельности [4]. Это показывает, что внедрение искусственного интеллекта в управление не сводится к подключению отдельной модели машинного обучения: требуется согласование алгоритма с системной моделью объекта и логикой принятия решения.

Для типовых управленческих ситуаций ранее предлагались архитектуры, объединяющие интеллектуальные и аналитические методы идентификации ситуации, категоризации и выбора сценария [5]. В современных работах журнала «Моделирование, оптимизация и информационные технологии» развиваются подходы к интеллектуальной поддержке управления распределенными организационными системами на основе структурного моделирования [6]. Вместе с тем остается методическая проблема разделения функций прогнозирования, оценки альтернатив и окончательного выбора: объединение этих этапов внутри единой непрозрачной модели затрудняет объяснение результата и локализацию ответственности.

Цель исследования состоит в разработке концептуально-математического подхода к применению искусственного интеллекта в системном анализе и поддержке принятия управленческих решений. Для достижения цели решаются следующие задачи: формализовать место ИИ в цикле системного анализа; предложить архитектуру интеллектуальной системы поддержки принятия решений; определить процедуру многокритериального выбора, отделяющую прогноз от управленческих предпочтений; исследовать чувствительность рейтинга альтернатив к изменению весов критериев; сформулировать требования к человеко-машинному взаимодействию и контролю рисков. Научная новизна работы заключается в совместном использовании функционально-слоистой архитектуры интеллектуальной поддержки и процедуры анализа устойчивости многокритериального рейтинга, что позволяет явно отделить

вычислительный прогноз от нормативного выбора и сохранить ответственность лица, принимающего решение.

Материалы и методы

Исследование имеет концептуально-модельный характер. В качестве объекта рассматривается организационно-техническая система, состояние которой описывается совокупностью наблюдаемых и скрытых параметров, а управление осуществляется посредством выбора одного из допустимых воздействий. Методическая схема включает четыре взаимосвязанных блока: системную формализацию объекта; декомпозицию функций интеллектуальной поддержки; многокритериальное ранжирование альтернатив; анализ чувствительности результата к изменению управленческих приоритетов. Такой подход соответствует направлению исследований человеко-ИИ систем в науках о принятии решений, где эффективность рассматривается как результат совместного проектирования технической и человеческой частей системы [7].

Управляемая система задается кортежем, включающим состояние, управляющие воздействия, наблюдаемые результаты, внешние факторы и отображение перехода между состояниями. Представление не предполагает, что функция перехода известна аналитически: ее отдельные компоненты могут идентифицироваться статистическими и машинно-обучаемыми моделями. Системная формализация нужна для фиксации границ объекта, определения того, какие переменные являются управляемыми, какие поступают из внешней среды и какие последствия должны контролироваться.

$$S = \langle X, U, Y, Z, F \rangle$$

Здесь X — множество состояний объекта; U — допустимые управляющие воздействия; Y — наблюдаемые результаты; Z — внешние возмущения и ограничения; F — совокупность зависимостей,

определяющих переходы и формирование результата. Искусственный интеллект может использоваться для оценки неизвестных компонентов F , прогнозирования Y при различных U и обнаружения изменений в структуре взаимосвязей. Однако формирование цели, границ системы и ограничений относится к системному уровню и не должно автоматически выводиться из исторических данных.

Второй элемент методики — функциональная декомпозиция интеллектуальной поддержки. Она основана на принципе разделения задач: методы анализа данных формируют признаки и оценки состояния; прогнозные модели оценивают последствия; многокритериальная процедура сопоставляет альтернативы; экспертный контур проверяет допустимость результата. Такой принцип особенно важен для объяснимого искусственного интеллекта: исследования показывают, что объяснимость должна проектироваться применительно к конкретной задаче и пользователю, а не рассматриваться как универсальное свойство модели [8].

Третий элемент — формальная модель выбора. Управленческая задача представляется через множество альтернатив, критерии, их веса и ограничения. В отличие от подхода, при котором модель ИИ непосредственно выдает «лучшее решение», в предлагаемой схеме алгоритм формирует прогнозные показатели и частные оценки, а агрегирование выполняется в явной многокритериальной процедуре.

$$D = \langle A, K, W, C \rangle$$

В выражении $A = \{A_1, \dots, A_m\}$ — множество допустимых альтернатив; $K = \{K_1, \dots, K_n\}$ — набор критериев; $W = \{w_1, \dots, w_n\}$ — веса критериев; C — ограничения, определяющие допустимую область решений. Для сопоставимости показатели нормируются. Для критерия, который необходимо максимизировать, используется отношение значения альтернативы к наилучшему значению в рассматриваемом множестве. Для

критерия затрат или риска применяется обратная нормализация. После этого определяется интегральная оценка.

Веса критериев задаются лицом, принимающим решение, либо экспертной группой. Допускается использование алгоритмических рекомендаций по весам на основе истории решений, однако окончательное утверждение W должно оставаться явной управленческой процедурой. Это связано с известной проблемой калибровки доверия к алгоритму: наличие дополнительной информации о точности ИИ и мотивация пользователя к проверке существенно влияют на качество человеко-машинного решения [9].

$$Q_j = \sum (w_i \cdot r_{ij}), i = 1, \dots, n; w_i \geq 0; \sum w_i = 1$$

Для анализа устойчивости используется набор сценариев изменения весов. Для каждой альтернативы вычисляется $Q_j(s)$, где s обозначает сценарий приоритетов. Дополнительно определяется диапазон вариации интегральной оценки. Чем выше этот диапазон и чем чаще меняется позиция альтернативы, тем сильнее решение зависит от субъективной настройки приоритетов. В задачах, где ошибка рекомендации имеет высокую цену, одной визуальной объяснимости недостаточно: эмпирические исследования показывают, что объяснение не всегда компенсирует негативный эффект ошибочного совета ИИ [10].

$$\Delta Q_j = \max Q_j(s) - \min Q_j(s)$$

Четвертый элемент методики — требования к объяснению и ответственности. При подготовке рекомендации система должна отображать ключевые исходные показатели, критерии, веса, ограничения и чувствительность результата. Подход учитывает, что объяснения алгоритмических решений ситуативны и могут сами становиться источником неправильной интерпретации [11]. Кроме того, даже при формальном сохранении «человека в контуре» остается проблема

атрибуции ответственности, если реальная практика превращает проверку рекомендации в формальность [12]. Поэтому экспертный контроль в предлагаемой архитектуре рассматривается как самостоятельный функциональный слой, а не как декларативное подтверждение результата.

Для проверки работоспособности подхода используется демонстрационный расчет выбора архитектуры интеллектуальной подсистемы. Сравниваются три альтернативы: A_1 — правилочная система; A_2 — преимущественно модель машинного обучения; A_3 — гибридная архитектура, в которой прогнозная модель дополняется правилами, ограничениями и экспертной проверкой. Оценивание проводится по пяти критериям: аналитическое качество, экономический эффект, объяснимость, интегрируемость и контроль риска. Значения по десятибалльной шкале имеют демонстрационный характер и не претендуют на эмпирическую репрезентативность; цель расчета — проверить поведение предложенной процедуры при различных наборах управленческих приоритетов.

Результаты

В результате функциональной декомпозиции сформирована шестислойная архитектура интеллектуальной системы поддержки принятия решений. Ее отличительной особенностью является отсутствие единственного автономного алгоритмического блока, который одновременно анализирует данные, прогнозирует состояние и формирует окончательное действие. Вместо этого каждый этап получает проверяемый вход и формирует интерпретируемый результат для следующего слоя. Архитектура представлена в таблице 1.

Таблица 1. Функциональные слои интеллектуальной системы поддержки принятия решений / Table 1. Functional layers of an intelligent decision support system

Слой / Layer	Содержание / Content	Функция ИИ / AI function	Результат / Output
1. Данные	Внутренние и внешние источники, телеметрия, документы	Очистка, извлечение признаков, обнаружение аномалий	Подготовленный набор данных
2. Системная модель	Цели, состояния, связи, ограничения, границы объекта	Идентификация ситуации, оценка скрытых состояний	Формализованная проблемная ситуация
3. Прогноз и сценарии	Варианты развития при альтернативных воздействиях	Прогнозирование, имитация, сценарный анализ	Оценки последствий альтернатив
4. Многокритериальная оценка	Критерии, веса, ограничения, риск	Расчет частных показателей и чувствительности	Рейтинг допустимых решений
5. Экспертный контроль	Проверка ограничений, объяснений и последствий	Формирование пояснений, поиск противоречий	Подтвержденное решение ЛПР
6. Обратная связь	Фактический эффект реализованного решения	Мониторинг ошибки, качества и дрейфа	Корректировка моделей и правил

Первый слой отвечает за качество исходной информации. В системном анализе ошибка на этапе данных опасна тем, что последующие вычислительные блоки могут усилить ее и представить как количественно точный результат. Поэтому до применения ИИ выполняются контроль полноты, временной согласованности, единиц измерения и происхождения данных. Второй слой связывает данные с моделью объекта: определяется, какое изменение относится к внутреннему состоянию, какое является внешним воздействием и какие ограничения должны быть выполнены независимо от прогнозного выигрыша.

Третий слой формирует прогнозы и сценарии. Здесь могут применяться модели временных рядов, классификация, регрессия, нейросетевые методы, имитационное моделирование или их комбинации. Четвертый слой превращает прогнозы в сопоставимые оценки альтернатив. На этом этапе принципиально отделяются фактические прогнозные величины от предпочтений ЛПР: веса критериев хранятся и изменяются независимо от прогнозной модели. Пятый слой обеспечивает экспертную проверку, включая анализ нарушений ограничений и объяснение факторов, повлиявших на результат. Шестой слой фиксирует фактический эффект и обеспечивает замыкание контура управления.

Предложенная архитектура позволяет локализовать источник ошибки. Если фактический результат расходится с прогнозом, проверяется третий слой; если прогноз был корректным, но выбранная альтернатива оказалась неудачной, анализируются критерии, веса и ограничения четвертого слоя; если решение принято вопреки предупреждению системы, проблема относится к процедуре экспертного контроля. Такая диагностируемость является преимуществом по сравнению с монолитной моделью, в которой невозможно отделить ошибку данных от ошибки прогнозирования или от неверно заданной целевой функции.

Таблица 2. Демонстрационная многокритериальная оценка альтернатив /

Table 2. Demonstrative multicriteria assessment of alternatives

Альтернатива / Alternative	Качество 0,30 / Quality	Экон. эффект 0,20 / Economic effect	Объяснимость 0,20 / Explainability	Интеграция 0,15 / Integration	Контроль риска 0,15 / Risk control	Q
A ₁ : правилковая / rule-based	6	9	9	8	9	7,95

А ₂ : ML- модель / ML model	9	6	5	7	5	6,7 0
А ₃ : гибридная / hybrid	8	8	8	9	8	8,1 5

При сбалансированном наборе весов наибольшую интегральную оценку получает гибридная альтернатива А₃ (Q = 8,15). Правилочная система А₁ уступает по аналитическому качеству, однако выигрывает по объяснимости и контролю риска. Модель машинного обучения А₂ имеет наивысшую частную оценку качества, но проигрывает по объяснимости и контролю риска. Полученный результат демонстрирует принцип системного выбора: максимизация локальной метрики прогнозной точности не гарантирует оптимальности решения на уровне всей управленческой системы.

Для оценки устойчивости проведены два дополнительных сценария. В сценарии повышенного риска вес критерия контроля риска увеличен до 0,35, а веса качества, экономического эффекта, объяснимости и интеграции заданы как 0,20; 0,15; 0,20 и 0,10 соответственно. В сценарии интенсивного развития приоритет смещен к аналитическому качеству и интеграции: веса равны 0,40; 0,15; 0,10; 0,25 и 0,10. Результаты приведены в таблице 3.

Таблица 3. Чувствительность рейтинга к изменению управленческих приоритетов / Table 3. Ranking sensitivity to changes in managerial priorities

Альтернатива / Alternative	Сбалансированный / Balanced	Повышенный риск / Risk- focused	Интенсивное развитие / Development- focused
А ₁	7,95	8,30	7,55
А ₂	6,70	6,15	7,25
А ₃	8,15	8,10	8,25

Чувствительность рейтинга показывает два практически значимых эффекта. Во-первых, гибридная система А₃ остается лидером в сбалансированном и ориентированном на развитие сценариях, то есть ее

преимущество не является следствием только одного конкретного набора весов. Во-вторых, при резком повышении значимости контроля риска лидером становится правилковая альтернатива A_1 . Следовательно, процедура корректно отражает изменение системной цели и не маскирует нормативные предпочтения внутри алгоритмической модели.

Для практической системы важен не только итоговый Q , но и причина его изменения. Поэтому в интерфейсе интеллектуальной СППР целесообразно отображать вклад каждого критерия, значения весов, статус ограничений и результат сценарного анализа. Если две альтернативы имеют близкие оценки, система должна сообщать о низкой устойчивости выбора, а не формировать категоричную рекомендацию. Такая логика поддерживает режим, в котором ИИ повышает ситуационную осведомленность и вычислительную эффективность, но не подменяет ответственность пользователя.

Обсуждение

Полученные результаты согласуются с современным пониманием интеллектуальной поддержки как социотехнической системы. В отличие от подходов, где искусственный интеллект оценивается исключительно по точности прогноза, предложенная архитектура рассматривает его как один из функциональных компонентов управленческого контура. Это соответствует работам по человеко-машинному сотрудничеству, в которых подчеркивается необходимость проектировать не только алгоритм, но и порядок обмена информацией, роли участников и механизмы проверки [2, 4].

Разделение прогноза и многокритериального выбора является принципиальным. Алгоритм машинного обучения может оценивать вероятность события, ожидаемый спрос, время выполнения операции или риск отказа, однако выбор действия требует сопоставления нескольких

последствий и учета ограничений, которые часто имеют нормативную или организационную природу. В предложенной модели такие параметры представлены явно через K , W и C . Это снижает риск ситуации, когда исторически сформированная метрика модели фактически превращается в скрытую целевую функцию всей системы.

Сравнение с современными исследованиями human-AI systems показывает, что проблема взаимодействия человека и алгоритма остается самостоятельной областью проектирования. Обзор направлений развития человеко-ИИ систем в decision sciences указывает на необходимость учитывать технические, поведенческие и организационные аспекты одновременно [7]. Для предлагаемой архитектуры это выражается в наличии экспертного слоя и обратной связи: рекомендация должна быть не только вычислена, но и включена в процедуру, позволяющую человеку понять основание выбора, проверить ограничения и зафиксировать принятое решение.

Отдельного внимания требует объяснимость. Современные работы по ХАИ связывают полезность объяснения с данными, методом и контекстом применения, а не только с наличием визуализации или перечня важных признаков [8]. В рассматриваемой СППР объяснение строится на системной структуре: показывается исходное состояние, прогнозируемое следствие, частные критерии, веса и чувствительность рейтинга. Такая форма объяснения ближе к управленческой задаче, поскольку позволяет ответить не только на вопрос «почему модель выдала прогноз», но и «почему именно эта альтернатива стала предпочтительной при данных приоритетах».

При этом объяснение не должно создавать иллюзию надежности. Экспериментальные исследования показывают, что пользователи могут переоценивать ошибочный совет ИИ, а наличие объяснения не всегда устраняет отрицательный эффект неверной рекомендации [10]. Поэтому в критичных задачах необходимо сочетать объяснение с калибровкой

доверия: показывать качество модели на сопоставимых данных, уровень неопределенности, случаи выхода за область применимости и основания, по которым рекомендация должна быть передана человеку без автоматического ранжирования.

Аналогичный вывод следует из исследований оптимизации человеко-машинного сотрудничества. Информация о точности ИИ и мотивация пользователя к осмысленной проверке способны влиять на итоговое качество решения [9]. В предлагаемой архитектуре этот вывод реализуется организационно: экспертный слой должен иметь регламент проверки, а не только кнопку подтверждения. Для сложных решений можно использовать принцип двухуровневого контроля, когда аналитик проверяет расчетные предпосылки, а руководитель утверждает управленческое действие с учетом целей и ограничений.

Важным ограничением ХАИ является контекстная зависимость объяснений. Исследования алгоритмического управления показывают, что объяснение может быть оспариваемым, неполным и неодинаково интерпретируемым различными группами пользователей [11]. Поэтому набор объясняющих элементов должен проектироваться под роль пользователя. Аналитику нужны диагностические показатели модели и ошибки на данных; руководителю — чувствительность решения, риск и ожидаемый эффект; специалисту по безопасности — ограничения, отказоустойчивость и границы применимости. Универсальный интерфейс, одинаково отображающий внутренние признаки модели всем пользователям, не обеспечивает системной прозрачности.

Проблема ответственности также требует явной архитектурной поддержки. Даже если формально последнее решение принимает человек, фактическая зависимость от рекомендации может создавать разрыв между правом принять решение и реальной возможностью его обосновать. В литературе этот эффект рассматривается как проблема атрибуции

ответственности в AI decision-support [12]. Предложенная схема частично уменьшает риск за счет журналирования исходных данных, версии модели, весов критериев, нарушенных ограничений, сформированной рекомендации и итогового действия ЛПР. Такой журнал создает трассируемость процесса и позволяет различать ошибку модели, ошибку параметризации и человеческое решение.

С точки зрения системного анализа существенным результатом является чувствительность рейтинга к управленческим приоритетам. Таблица 3 показывает, что выбор архитектуры не может быть окончательно определен только по средним техническим характеристикам. При росте цены ошибки прозрачная правилковая система становится предпочтительной, хотя ее аналитическая точность ниже. В сценарии развития преимущество гибридной альтернативы увеличивается. Следовательно, проектирование интеллектуальной СППР должно начинаться с формализации системной цели и допустимого риска, а не с выбора конкретной модели ИИ.

Практическое внедрение предложенного подхода целесообразно выполнять поэтапно. На первом этапе фиксируются границы системы, цель, показатели результата и критические ограничения. На втором проводится аудит данных и формируется базовый способ принятия решения, с которым будет сравниваться интеллектуальная подсистема. На третьем реализуется прогнозный модуль и проверяется его устойчивость на исторических и новых данных. На четвертом добавляется многокритериальный слой с явным хранением весов и правил нормирования. Только после этого формируется человеко-машинный интерфейс и процедура эксплуатационного контроля.

Для пилотного проекта рекомендуется выбирать повторяющуюся управленческую задачу с измеримым результатом. К техническим метрикам модели необходимо добавить системные показатели: время подготовки

решения, частоту критических ошибок, число отклоненных рекомендаций, устойчивость рейтинга альтернатив, экономический эффект и долю случаев, когда модель работала вне области применимости. Если пользователь систематически отклоняет рекомендации, это следует рассматривать не как «ошибку человека», а как диагностический сигнал о несоответствии модели, критериев или интерфейса реальной задаче.

Ограничения исследования связаны с концептуальным характером предложенной архитектуры и демонстрационным характером числового примера. Значения критериев не получены в натурном эксперименте и используются для проверки поведения процедуры. Для эмпирической валидации необходимы отраслевые данные, экспертное определение шкал, сравнение нескольких способов нормирования и экспериментальная оценка качества решений с ИИ и без него. Отдельным направлением дальнейшей работы является разработка методов автоматического обнаружения дрейфа системной ситуации, при котором ранее установленные веса и ограничения перестают соответствовать текущей среде.

Тем не менее предложенный подход имеет методическое значение для направления «Системный анализ и управление». Он связывает задачи идентификации, прогнозирования и интеллектуального анализа данных с классическими функциями системного анализа: определением границ объекта, целей, ограничений, критериев и обратных связей. Это позволяет рассматривать искусственный интеллект не как автономного субъекта управления, а как инструмент расширения аналитических возможностей внутри контролируемого контура принятия решения.

Заключение

Разработан концептуально-математический подход к включению искусственного интеллекта в системный анализ и поддержку принятия управленческих решений. В отличие от монолитной схемы, где алгоритм

непосредственно формирует решение, предложена шестислойная архитектура, разделяющая подготовку данных, системное моделирование, прогнозирование, многокритериальную оценку, экспертный контроль и обратную связь.

Формальная процедура выбора основана на явном представлении альтернатив, критериев, весов и ограничений. Прогнозные возможности ИИ используются для получения частных оценок, тогда как окончательное ранжирование зависит от заданных управленческих приоритетов. Анализ чувствительности позволяет выявлять решения, рейтинг которых существенно меняется при вариации весов, и тем самым предотвращать необоснованно категоричные рекомендации.

Демонстрационный расчет показал преимущество гибридной архитектуры при сбалансированном и ориентированном на развитие наборах критериев. При существенном повышении веса контроля риска предпочтительной становится правилковая система, что подтверждает необходимость явного учета системной цели и цены ошибки. Полученный результат показывает, что наиболее точная прогнозная модель не обязательно является наиболее предпочтительным элементом управленческой системы.

Практическая реализация подхода требует журналирования данных и параметров решения, отображения вклада критериев, контроля области применимости модели и регламентированной экспертной проверки. Перспективы дальнейших исследований связаны с апробацией архитектуры на реальных организационно-технических задачах, сравнением методов определения весов, разработкой индикаторов дрейфа системной ситуации и оценкой качества человеко-машинных решений в условиях неопределенности.

Сведения об использовании генеративного искусственного интеллекта

При подготовке рабочего варианта рукописи использован генеративный инструмент для создания чернового текста, языкового редактирования и структурирования материала. Перед подачей автор должен самостоятельно проверить и переработать научные положения, расчеты, ссылки, интерпретацию результатов и окончательные выводы.¹

СПИСОК ИСТОЧНИКОВ / REFERENCES

1. Грушо А.А., Забежайло М.И., Зацаринный А.А. Обеспечение ситуационной осведомленности в системах поддержки управленческих решений на основе технологий искусственного интеллекта. Искусственный интеллект и принятие решений. 2026;(1):3–14. <https://doi.org/10.14357/20718594260101>
Grusho A.A., Zabezhailo M.I., Zatsarinnyi A.A. Maintaining situational awareness in management decision support systems based on artificial intelligence technologies. Artificial Intelligence and Decision Making. 2026;(1):3–14. (In Russ.). <https://doi.org/10.14357/20718594260101>
2. Смирнов А.В., Левашова Т.В. Паттерны человеко-машинного сотрудничества в системах поддержки принятия решений. Искусственный интеллект и принятие решений. 2024;(2):3–17. <https://doi.org/10.14357/20718594240201>
Smirnov A.V., Levashova T.V. Patterns of human-machine collaboration in decision support systems. Artificial Intelligence and Decision Making. 2024;(2):3–17. (In Russ.). <https://doi.org/10.14357/20718594240201>
3. Дородных Н.О., Николайчук О.А., Котлов Ю.В., Юрин А.Ю. Создание диагностических систем поддержки принятия решений с использованием таблиц решений. Искусственный интеллект и принятие решений. 2025;(3):15–29.

¹

<https://doi.org/10.14357/20718594250302>

Dorodnykh N.O., Nikolaychuk O.A., Kotlov Yu.V., Yurin A.Yu. Creation of diagnostic decision support systems using decision tables. Artificial Intelligence and Decision Making. 2025;(3):15–29. (In Russ.). <https://doi.org/10.14357/20718594250302>

4. Смирнов А.В., Левашова Т.В. Онтология паттернов человеко-машинного сотрудничества для поддержки принятия решений. Онтология проектирования. 2024;14(3):421–439. <https://doi.org/10.18287/2223-9537-2024-14-3-421-439>

Smirnov A.V., Levashova T.V. Ontology of human-machine collaboration patterns for decision support. Ontology of Designing. 2024;14(3):421–439. (In Russ.). <https://doi.org/10.18287/2223-9537-2024-14-3-421-439>

5. Антонов В.В., Конев К.А., Куликов Г.Г. Трансформация модели системы поддержки принятия решений для типовых ситуаций с применением интеллектуальных и аналитических методов. Вестник ЮУрГУ. Серия «Компьютерные технологии, управление, радиоэлектроника». 2021;21(3):14–25. <https://doi.org/10.14529/ctcr210302>

Antonov V.V., Konev K.A., Kulikov G.G. Transformation of the decision support system model for typical situations using intelligent and analytical methods. Bulletin of the South Ural State University. Ser. Computer Technologies, Automatic Control, Radio Electronics. 2021;21(3):14–25. (In Russ.). <https://doi.org/10.14529/ctcr210302>

6. Максим А.Д. Формирование интеллектуальной поддержки управленческих решений в территориально распределенной организационной системе с персонализированными субъектами на основе структурного моделирования. Моделирование, оптимизация и информационные технологии. 2026;14(5). <https://doi.org/10.26102/2310->

7. Storey V.C., Hevner A.R., Yoon V.Y. The design of human-artificial intelligence systems in decision sciences: A look back and directions forward. *Decision Support Systems*. 2024;182:114230. <https://doi.org/10.1016/j.dss.2024.114230>
8. Coussement K., Abedin M.Z., Kraus M., Maldonado S., Topuz K. Explainable AI for enhanced decision-making. *Decision Support Systems*. 2024;184:114276. <https://doi.org/10.1016/j.dss.2024.114276>
9. Eisbach S., Langer M., Hertel G. Optimizing human-AI collaboration: Effects of motivation and accuracy information in AI-supported decision-making. *Computers in Human Behavior: Artificial Humans*. 2023;1(2):100015. <https://doi.org/10.1016/j.chbah.2023.100015>
10. Cecil J., Lermer E., Hudecek M.F.C., et al. Explainability does not mitigate the negative impact of incorrect AI advice in a personnel selection task. *Scientific Reports*. 2024;14:9736. <https://doi.org/10.1038/s41598-024-60220-5>
11. de Bruijn H., Warnier M., Janssen M. The perils and pitfalls of explainable AI: Strategies for explaining algorithmic decision-making. *Government Information Quarterly*. 2022;39(2):101666. <https://doi.org/10.1016/j.giq.2021.101666>
12. Zeiser J. Owning Decisions: AI Decision-Support and the Attributability-Gap. *Science and Engineering Ethics*. 2024;30:27. <https://doi.org/10.1007/s11948-024-00485-1>

ИНФОРМАЦИЯ ОБ АВТОРАХ / INFORMATION ABOUT THE AUTHORS

Мазеин Кирилл Сергеевич	Mazein Kirill Sergeevich
Студент	Student
ФГБОУ ВО «Санкт-Петербургский государственный технологический институт (технический университет)»	Federal State Budgetary Educational Institution of Higher Education “Saint Petersburg State Technological Institute (Technical University)”
Санкт-Петербург	Saint-Petersburg
Российская Федерация	Russian Federation
ORCID: 0009-0005-3792-1692	ORCID: 0009-0005-3792-1692