

Формализованная модель отбора процессов предприятия для внедрения технологий искусственного интеллекта

Бондарь Д. П.

Аннотация

Эмпирические исследования фиксируют, что большинство проектов внедрения искусственного интеллекта (ИИ) на предприятиях не достигает продуктивной эксплуатации; доминируют организационные причины: неверный выбор задачи, отсутствие данных и метрик успеха, определённых до старта. Существующие инструменты оценки либо созданы для статистики занятости и измеряют пригодность задач к автоматизации в масштабе экономики, а не отдельного предприятия (рубрика пригодности задач для машинного обучения — SML — и семейство мер подверженности задач воздействию ИИ — exposure), либо диагностируют организационную готовность, не выходя ни на отбор конкретных инициатив, ни на требования к ведению отобранного проекта, либо требуют многофазных процедур с фасилитаторами (методы многокритериального анализа решений). Инструмент, позволяющий владельцу процесса за одну рабочую сессию обоснованно решить, передавать ли задачу ИИ, в рассмотренной литературе не обнаружен. Цель исследования — разработать и обосновать формализованную модель отбора процессов для внедрения ИИ. Исследование выполнено в парадигме design science research. Предложенная модель включает четырнадцать атрибутов процесса по шкале 0/1/2, четыре некомпенсаторных стоп-условия, индекс осуществимости с порогами решения, индекс приоритета с экономической формулой эффекта и правила безопасной реализации. Состав атрибутов не назначен экспертно: он выведен из критериев рубрики SML, факторов организационной готовности и задокументированных причин неудач ИИ-проектов, а стоп-условия соответствуют четырём самым частым из них. Проведено попутное сопоставление с рубрикой SML и ретроспективное кодирование 14 опубликованных кейсов (7 неудач, 7 успехов) по протоколу против ретроспективного искажения. Модель различает группы: у шести из семи неудач зафиксированы сработавшие стоп-условия и меньшие значения индекса; у успехов сработавших стоп-условий нет. Результат интерпретируется как ретроспективная согласованность с известными исходами; проспективная валидация — направление дальнейших исследований.

Ключевые слова: искусственный интеллект; технологии искусственного интеллекта; бизнес-процессы; управление проектами; поддержка принятия решений; машинное обучение; цифровая трансформация; design science research.

Введение

Разрыв между технической пригодностью задачи для искусственного интеллекта и фактическим успехом внедрения стал устойчивым эмпирическим фактом. По оценкам, приводимым в отчёте RAND, более 80% ИИ-проектов не достигают целей — вдвое больше, чем у ИТ-проектов без ИИ-компоненты; пять корневых причин носят

преимущественно организационный характер: неверная постановка задачи, отсутствие данных, ориентация на технологию вместо проблемы, недостаток инфраструктуры и выбор задач, превышающих возможности технологии [Ryseff et al., 2024]. Отраслевые оценки доли неудач колеблются от 30% (прогноз по отказам от проектов генеративного ИИ [Gartner, 2024]) до 95% (доля пилотов без измеримого финансового эффекта [MIT NANDA, 2025]). Эти цифры получены разными методами на разных определениях «неудачи» — от недостижения окупаемости до полной остановки, — и сама их несопоставимость показательна: общепринятого критерия успеха внедрения ИИ не сложилось. Процедуры предпроектного обоснования при этом существуют — технико-экономическое обоснование, стадийные фильтры портфельного управления, — но специфику ИИ-проектов они не учитывают: зависимость результата от данных, вероятностное качество, смещение поведения системы после запуска.

Дело не только в исполнении: сам способ выбирать, что автоматизировать, оценивает не то. Там, где прогнозы автоматизации можно проверить на длинной дистанции — на рынке труда, — предсказания «что ИИ может» разошлись с тем, что внедрили на деле (ретроспективные проверки — в разделе 1.1). Для предприятия вывод прямой: отбор задач по одной лишь технической пригодности воспроизведёт этот разрыв, потому что успех внедрения зависит ещё от данных, интеграции, экономики и организационных условий конкретного предприятия [Frank et al., 2019].

Существующие инструменты оценки не закрывают этот разрыв. Рубрика пригодности задач для машинного обучения (*suitability for machine learning*, SML) [Brynjolfsson, Mitchell, 2017] и семейство мер подверженности задач воздействию ИИ (*exposure*) созданы для экономической статистики рынков труда и, по признанию самих авторов, не учитывают экономические, организационные и правовые факторы. Фреймворки организационной готовности к ИИ диагностируют организацию в целом, но не дают ни процедуры отбора конкретной инициативы, ни требований к ведению отобранного проекта — базовых замеров, порогов качества, условий расширения автономии; методы многокритериального анализа решений требуют экспертных панелей и фасилитации. Инструмента, которым руководитель, отвечающий за конкретный повторяющийся процесс, мог бы за одну рабочую сессию получить обоснованный ответ — передавать ли эту задачу ИИ и с каким приоритетом, — в рассмотренной литературе не обнаружено.

Цель настоящего исследования — разработать и обосновать формализованную модель отбора процессов предприятия для внедрения ИИ, соединяющую в одном компактном инструменте техническую осуществимость, экономику эффекта, организационную готовность и условия безопасной реализации. Исследование выполнено в парадигме *design science research* (DSR): проектируемый артефакт — метод отбора; его обоснование строится на выводимости каждого элемента модели из опубликованных эмпирических результатов, а оценка — на ретроспективном кодировании опубликованных

кейсов внедрений с задокументированными исходами. Поэтому изложение начинается с обзора: раздел 1 проверяет заявленный пробел в четырёх пластах литературы — мерах пригодности задач, эмпирике неудач и организационной готовности, методах отбора проектов и русскоязычных работах — и одновременно предоставляет материал, из которого выводятся элементы модели.

1. Обзор литературы

1.1. Оценка пригодности задач для ИИ: рубрика SML и меры подверженности

Линия количественной оценки автоматизируемости прошла путь от уровня профессий (47% занятости США в зоне высокого риска у [Frey, Osborne, 2017]) к уровню задач: перенос анализа на задачи снижает оценку до 9–14%, поскольку внутри одной профессии сосуществуют задачи разной автоматизируемости [Arntz et al., 2016; Nedelkoska, Quintini, 2018]. Этот сдвиг — от роли к задаче — отправная точка и для настоящего исследования: объект отбора у предлагаемой модели не «профессия» и не «направление цифровизации», а отдельный повторяющийся процесс.

Следующий шаг сделали Э. Бриньольфссон и Т. Митчелл, предложив оценивать пригодность задач по восьми ключевым критериям рубрики SML — от чётко определённых входов и выходов и существующих наборов данных «вход–выход» до толерантности к ошибкам и стабильности феномена [Brynjolfsson, Mitchell, 2017]. Развёрнутая 21-пунктная рубрика (полный текст пунктов воспроизведён в [Septiandri et al., 2024]) была применена в расширенной 23-пунктной версии к 2 069 рабочим активностям базы O*NET силами привлечённых разметчиков (краудворкеров) [Brynjolfsson et al., 2018]. Главные выводы: задачи с высокой пригодностью присутствуют почти во всех профессиях, но ни одна профессия не состоит из них целиком, поэтому раскрытие потенциала машинного обучения требует перепроектирования рабочих процессов.

Принципиальны ограничения, признанные самими авторами: рубрика оценивает только техническую осуществимость и не учитывает экономические, организационные и правовые факторы [Brynjolfsson et al., 2018]. Ретроспективные проверки этой линии отрезвляют: занятость в профессиях, отнесённых к зоне высокого риска автоматизации, за 2012–2019 гг. не сократилась, а в среднем выросла [Georgieff, Milanez, 2021]; корреляция расчётного потенциала автоматизации с фактическим распространением технологий — лишь 0,36 [Meindl et al., 2021]. Крауд-оценка абстрактных описаний задач даёт узкое распределение оценок и слабую предсказательную силу: при сопоставлении с фактическими вакансиями агрегированная оценка по рубрике уступает альтернативным индексам [Acemoglu et al., 2022]. Для больших языковых моделей потребовалась уже новая рубрика [Eloundou et al., 2024] — свидетельство того, что привязанные к поколению технологии критерии устаревают. Ту же дыру не закрывает и остальное семейство: меры подверженности через способности [Felten et al., 2021], патенты [Webb, 2020],

когнитивные способности и ИИ-бенчмарки [Tolan et al., 2021] — макроинструменты для статистики рынков труда, ни один не предназначен для решения о судьбе одного корпоративного процесса. Для настоящего исследования из этой линии важен поэтому не прогноз, а состав критериев — какие свойства делают задачу пригодной для машинного обучения: именно их предлагаемая модель наследует и переводит в проверку конкретного процесса (раздел 4).

1.2. Причины неудач ИИ-проектов и организационная готовность

Эмпирика неудач сходится на доминировании организационно-человеческих факторов. Помимо упомянутых корневых причин RAND, интервью с экспертами выделяют нереалистичные ожидания, неудачный выбор задачи и нехватку ключевых ресурсов, включая размеченные данные [Westenberger et al., 2022]. Инженерная литература документирует, что работоспособный прототип не равен продуктивной системе: скрытый технический долг систем машинного обучения [Sculley et al., 2015] и трудности этапа развёртывания [Paleyes et al., 2022] относятся к интеграции, данным и сопровождению, а не к качеству модели.

Второй пласт — фреймворки организационной готовности к ИИ. Они отвечают на вопрос, какие условия должны сложиться в организации, чтобы внедрение имело шансы: от стратегии и культуры до данных и инфраструктуры. Таких рамок много, и устроены они по-разному — пять категорий факторов [Jöhnk et al., 2021], четыре измерения [Holmström, 2022], двадцать три фактора в систематическом обзоре [Ali, Khan, 2025], — но общее у них одно: они оценивают готовность организации в целом и не дают процедуры отбора конкретной инициативы. Диагноз «организация готова» не отвечает на вопрос, какой процесс брать первым. При этом практический вывод консалтинговой эмпирики устойчив: инкрементальные проекты с измеримым эффектом успешнее амбициозных «прорывных» программ [Davenport, Ronanki, 2018].

1.3. Отбор и приоритизация проектов: от портфельных рамок к инструментам уровня задачи

Классическая рамка портфельного отбора разделяет предварительный отсев (screening) и балльную оценку прошедших кандидатов [Archer, Ghasemzadeh, 1999]. Эта рамка — концептуальный предок связки «стоп-условия → балльная оценка (скоринг)», однако применяется она на уровне портфельного офиса. Основное направление многокритериального анализа решений (multi-criteria decision making, MCDM) — метод анализа иерархий (АНР), TOPSIS и их гибриды — требует матриц парных сравнений, контроля согласованности и экспертных панелей. Стоимость применения самого метода не выступает критерием проектирования ни в одной из рассмотренных работ; в контексте ИИ эти методы применяются к выбору технологий и ранжированию барьеров, но не к отбору бизнес-задачи её владельцем.

Специализированные фреймворки отбора ИИ-кейсов появились в информационных системах недавно. Метод целенаправленной разработки ИИ-кейсов [Hofmann et al., 2020] и подход к управлению портфелем ИИ-кейсов [Bodendorf, 2025] представляют собой многофазные процедурные модели с групповыми экспертными сессиями (воркшопами). У первого критерии оценки не операционализированы и вырабатываются самой организацией; у второго ценность и осуществимость агрегируются компенсаторно, а финансовая ценность оценивается экспертным баллом, не расчётной формулой. Обзорная работа [Kakuschke et al., 2025] прямо констатирует отсутствие целостных прикладных инструментов оценки кейсов создания ценности из данных. Наконец, канва ИИ-проекта Агравала, Ганса и Голдфарба [Agrawal et al., 2018] — прецедент одностороннего инструмента с экономической логикой «ценность = экономика предсказания», но без скоринга, сравнения кандидатов и стоп-условий.

1.4. Русскоязычный контекст

В русскоязычной литературе задача ранжирования ИТ-проектов ставилась неоднократно [Грекул и др., 2019; Марон, Файнбург, 2025; Долганова, Деева, 2019]. Обширный пласт работ по цифровой зрелости предприятий [Гилева, 2019; Краковская и др., 2024] использует зрелость как диагностику состояния, однако связка «оценка зрелости → приоритизация конкретных ИИ-инициатив» не операционализирована. Литература по внедрению ИИ в промышленности сама фиксирует методический дефицит: «недостатки методологии цифровизации» называются системной проблемой [Комаров и др., 2024], отдельно констатируется отсутствие надёжных методик оценки эффектов [Белова, Бадина, 2025]. Ближайшая публикация — модель зрелости и факторы успеха внедрения ИИ-агентов [Лобозкий и др., 2026] — рассматривает стадию после решения о внедрении; предлагаемая модель закрывает предшествующий этап: какие процессы и в каком порядке отдавать ИИ. Критерии безопасности и рисков ИИ в найденных русскоязычных методиках отбора отсутствуют либо упоминаются как внешний барьер.

Таким образом, во всех четырёх пластах литературы подтверждается один и тот же пробел: отсутствует компактный формализованный инструмент отбора уровня отдельного процесса, соединяющий техническую осуществимость, экономику, организационные условия и требования безопасной реализации.

2. Методология исследования

Подтверждённый обзором пробел — отсутствующий инструмент — является задачей проектирования, а не эмпирической проверки гипотез, поэтому исследование выполнено в парадигме design science research [Hevner et al., 2004]: его результат — спроектированный артефакт, и парадигма даёт явные критерии, по которым такой результат обосновывается и проверяется; этим критериям подчинена дальнейшая структура статьи. Проектируемый артефакт по типологии [March, Smith, 1995] — метод: формализованная процедура отбора

процессов с моделью скоринга. Процесс исследования следует шести шагам методологии DSRM (design science research methodology) [Peppers et al., 2007] с проблемно-ориентированным входом: от идентификации проблемы (введение, раздел 1) через цели решения и разработку артефакта (разделы 2–3) к демонстрации и оценке (раздел 5); коммуникация структурирована по схеме публикации DSR-результатов [Gregor, Hevner, 2013].

Цели решения сформулированы явно, поскольку именно они определяют проектные решения модели: (а) заполняемость владельцем процесса за одну рабочую сессию без фасилитатора и специальной подготовки; (б) некомпенсаторность критических условий — высокая ценность не должна перекрывать отсутствие условий, при которых внедрение осуществимо; (в) явная экономика — эффект должен вычисляться формулой из измеримых величин, а не назначаться экспертным баллом; (г) отбор не заканчивается решением «брать»: вместе с ним фиксируются условия исполнения — замер до старта, порог приёма и правила расширения автономии.

Строгость (rigor cycle по [Hevner, 2007]) обеспечивается выводимостью каждого атрибута модели из опубликованных результатов: критериев SML, факторов готовности и задокументированных причин неудач (таблица 1). Релевантность (relevance cycle) обеспечена происхождением модели из практики внедрений на промышленных предприятиях. Уровень вклада по классификации [Gregor, Hevner, 2013] — зарождающаяся теория проектирования (уровень 2: конструкты, принципы, метод) в квадранте «улучшение» (известная проблема — новое решение). Для читателя эта классификация калибрует притязания работы: заявляются метод и принципы его построения, а не универсальная теория, и результаты раздела 5 следует оценивать по меркам именно этого уровня.

3. Формализованная модель отбора

Раздел предъявляет сам артефакт; каждый его элемент закрывает одну из целей раздела 2: атрибуты и анкета (3.1) — заполняемость владельцем процесса (а), стоп-условия (3.2) — некомпенсаторность (б), индексы осуществимости и приоритета (3.3–3.4) — явную экономику и правило решения (в), правила безопасной реализации (3.5) — условия исполнения (г); атрибуты реализуемости и правило быстрых побед (3.4) отделяют неопределённо реализуемые кандидаты от отработанных; направления инвентаризации (3.6) задают вход модели.

3.1. Четырнадцать атрибутов процесса

Модель оценивает отдельный повторяющийся процесс предприятия по четырнадцати атрибутам по порядковой шкале 0/1/2 с содержательными якорями; каждый атрибут имеет обоснование в литературе (таблица 1).

Таблица 1. Атрибуты модели и их основания

№	Атрибут (значение 2 / 1 / 0)	Основание
1	Однотипность выполнения: каждое выполнение похоже на предыдущее / несколько типовых вариантов / каждый раз по-разному	SML: повторяемость и схожесть экземпляров [Brynjolfsson, Mitchell, 2017]
2	Формализуемость правил: письменная инструкция / правила в головах экспертов / правил нет	SML: описываемость правилами, чёткие входы и выходы
3	Класс работы: сопоставление, извлечение, классификация, расчёт, контроль по изображению, генерация текста / смешанная / физические действия, переговоры, ответственные решения	SML: соответствие класса задачи возможностям технологии
4	Доля исключений («не по правилам»): <10% / 10–30% / >30%	SML: отсутствие длинных цепочек рассуждений на фоновых знаниях
5	Оцифрованность входа: полностью в системах / частично / бумага, голос	SML: машиночитаемость входов; причина неудач «нет данных» [Ryseff et al., 2024]
6	История выполнения с правильными результатами: годы истории / немного примеров / истории нет	SML: существующие наборы «вход–выход»
7	Контрольные случаи для замера (стоп): 100+ случаев с известным ответом / десятки / собрать нельзя	Отсутствие метрик до старта — сквозной диагноз неудач [Ryseff et al., 2024; Westenberger et al., 2022]
8	Точка встраивания (стоп): система/API есть / нужна доработка / встроиться некуда	Интеграция в рабочий поток — фактор развёртывания [Paleyes et al., 2022]
9	Цена ошибки: низкая, ошибка ловится / средняя, обратима / высокая, необратимый ущерб	SML: толерантность к ошибкам; экономическая компонента модели

№	Атрибут (значение 2 / 1 / 0)	Основание
10	Владелец-спонсор (стоп): активно заинтересован / нейтрален / владельца нет	Готовность: поддержка руководства и владение инициативой [Jöhnk et al., 2021; Ali, Khan, 2025]
11	Эксперты на верификацию: выделены регулярно / от случая к случаю / недоступны	Готовность: доступ к предметной экспертизе [Jöhnk et al., 2021; Ali, Khan, 2025]
12	Регламент допускает изменение исполнения (стоп): да / через согласование / участок неприкасаем	Перепроектирование работ как условие эффекта [Brynjolfsson et al., 2018]
13	Определённость технической реализации: отработанные решения / нужна адаптация / способ неясен	Организационная причина неудач — задача не по силам технологии [Ryseff et al., 2024]
14	Стабильность правил процесса: стабильны годами / предсказуемый цикл / меняются быстрее выхода на порог	SML: изучаемый феномен не должен быстро меняться во времени [Brynjolfsson, Mitchell, 2017]

Отдельного пояснения требует атрибут 3: его шкала — не градация качества, а состав работы. Значение 2 ставится, когда процесс целиком состоит из перечисленных информационных операций; 1 — когда они перемешаны с ручными действиями (расчёт может выполнять ИИ, а загрузку материалов — только человек); 0 — когда ядро процесса составляют физические действия, переговоры или решения с персональной ответственностью. Сам перечень информационных операций — это классы работ, надёжно выполняемые текущим поколением технологий, и каждый класс замыкается на определённый класс технических решений при выборе инструмента; отсюда технологическая зависимость этого атрибута, оговорённая в разделе 6. Атрибут 4 считает исключения — выполнения процесса, которые не укладываются в типовые правила и требуют нестандартного разбора человеком: нештатный документ, редкий случай, спорная ситуация; доля берётся от общего числа выполнений. Чем она выше, тем меньшую часть работы ИИ может взять на себя и тем дороже обходится сопровождение остатка — в уроках таблицы 2 именно обработка исключений персоналом обнуляла экономию. Атрибут 6 требует не просто журнала выполнения, а истории с известным верным итогом: для каждого прошлого выполнения зафиксировано, каким оказался правильный результат — фактическое время прокатки, подтверждённые расхождения сверки, проверенные юристом поля договора. Такие пары «вход — верный выход» — материал, на котором

система обучается и настраивается; атрибут 7 из того же материала требует отдельного: проверочного набора для замера качества. Атрибут 11 фиксирует организационное обязательство, а не наличие специалистов в штате: выделены ли предметным экспертам регулярные часы на верификацию знаний системы перед запуском, подтверждение контрольных случаев и разбор ошибок пилота. Типичная гибель пилота выглядит именно так: эксперты в организации есть, но время им не выделено — ошибки не разбираются, и качество застревает ниже порога. Атрибуты 1–4 (технический блок) операционализируют критерии пригодности; атрибуты 5–7 (блок данных) — требования к данным и к возможности измерить результат; атрибуты 8, 10–12 (организационный блок) — условия интеграции и владения; атрибуты 13–14 (блок реализуемости) — определённость технического решения и стабильность правил процесса; атрибут 9 связывает пригодность с экономикой. Определения атрибутов сформулированы технологически нейтрально — как свойства процесса и организации, а не возможностей конкретного поколения моделей: это защищает модель от устаревания, постигшего привязанные к технологии рубрики (см. [Eloundou et al., 2024]).

3.2. Стоп-условия

Четыре атрибута (7, 8, 10, 12) являются стоп-условиями: значение 0 по любому из них останавливает рассмотрение кандидата независимо от остальных оценок. Атрибут 13 действует иначе: значение 0 (способ реализации неясен) не отбраковывает кандидата, а отправляет его в короткую техническую разведку — проверку прототипом — до ранжирования, чтобы неопределённо реализуемый кандидат не конкурировал с отработанными на равных. Атрибут 14 — обычный балльный: низкая стабильность правил снижает индекс, но не останавливает отбор. Некомпенсаторная логика отличает модель и от SML, где итог задачи — среднее оценок по пунктам рубрики и провал по одному пункту перекрывается высокими баллами по другим, и от взвешенных сумм MCDM-методов: невозможность измерить результат, отсутствие точки встраивания, отсутствие владельца или неприкасаемость регламента не компенсируются никакой ожидаемой ценностью. Стоп-условия наследуют стадию предварительного отсева портфельной рамки [Archer, Ghasemzadeh, 1999], свёрнутую в инструмент уровня одного кандидата, и операционализируют четыре наиболее часто документируемые организационные причины неудач [Ryseff et al., 2024; Westenberger et al., 2022].

3.3. Индекс осуществимости

Индекс осуществимости вычисляется как $F = (\sum a_i / 28) \times 100$, где a_i — значения четырнадцати атрибутов; $F \in [0; 100]$. Пороги делят шкалу на три зоны: $F \geq 70$ — кандидат идёт в ранжирование; $45 \leq F < 70$ — сначала закрыть слабые атрибуты (типовые действия: оцифровка входа, накопление контрольных случаев, назначение владельца, согласование изменения регламента) и пересчитать; $F < 45$ — не рассматривать сейчас. Равные веса и аддитивная свёртка — осознанное проектное решение в пользу цели (а) из раздела 2:

заполнение владельцем процесса без специалиста по MCDM. На фоне отсутствия эмпирически валидированных весов во всём рассмотренном пласте (раздел 1.3) усложнение агрегирования не даёт доказанного преимущества, но лишает инструмент применимости без специалиста по MCDM; пороги 70/45 калиброваны экспертно на стадии формирующей оценки, что зафиксировано как ограничение (раздел 6).

3.4. Индекс приоритета

Прошедшие фильтр кандидаты ранжируются индексом $\Pi = E \times (F/100) / C$, где E — годовой экономический эффект, вычисляемый по формуле $E = N \times t \times c \times k + L \times k'$ (N — частота выполнения в год; t — трудоёмкость одного выполнения; c — стоимость часа исполнителя; L — годовые потери от ошибок процесса; k, k' — достижимые доли улучшения), а C — класс стоимости внедрения (порядковая шкала S/M/L). Π интерпретируется только порядково — как ранжирующий показатель для сравнения кандидатов внутри одного перечня, а не абсолютная оценка рентабельности: множитель $F/100$ — эвристический дисконт эффекта на осуществимость, а не оценка вероятности успеха. Классам стоимости присваиваются фиксированные числовые значения ($S = 1, M = 3, L = 9$ — по типовым порядкам затрат); итоговый порядок зависит от выбранного кодирования, поэтому кодировка фиксируется до начала оценивания и едина для всего перечня. Доли улучшения k и k' — наиболее субъективные параметры формулы: при отсутствии аналогов они оцениваются консервативно и уточняются по результатам пилота; правило базовой линии (раздел 3.5) превращает их из ожидания в проверяемую величину. Формула операционализирует экономическую логику «ценность = экономика предсказания» [Agrawal et al., 2018], которой не содержит в расчётном виде ни один из рассмотренных инструментов отбора. Ранжирование по Π дополняют два правила, устраняющих его известную слабость — недооценку риска нереализуемости и ценности быстрой победы. Правило быстрых побед: кандидат с высокой осуществимостью, отработанной реализацией (атрибут 13 = 2) и низкой стоимостью внедрения берётся параллельно основному вектору, вне общего ранжирования, — он дешёв, но приносит доверие организации и обкатанную инфраструктуру для следующих векторов. Вторая ось выбора — оценка времени до порога приёмки (в формулу не входит): частый процесс с коротким циклом обратной связи выходит на порог быстрее редкого. Без этих правил редкий, но простой и надёжный процесс проигрывал бы в ранжировании частому процессу с неопределённым способом реализации, хотя довести второй до эффекта существенно менее вероятно.

3.5. Условия безопасной реализации

Модель связывает отбор с реализацией двумя правилами. Правило базовой линии: до старта пилота фиксируется измеренное состояние процесса; для систем класса «вопрос–ответ» — прямым тестом на выборке из ста и более вопросов с заранее известными ответами, доля верных ответов образует и точку отсчёта, и порог приёмки. Правило

заработанной автономии: сокращение ручного труда и изменение регламента допускаются только после устойчивого удержания метрики качества выше согласованного порога; автономия системы расширяется по мере накопленных подтверждений качества. Первое правило — прямой ответ на диагноз «нет метрик до старта», второе согласуется с выводами о постепенном расширении автономии и человеке в контуре [Лобозкий и др., 2026] и служит встроенным механизмом адаптации модели к росту возможностей технологий.

3.6. Направления инвентаризации

Вход модели — перечень кандидатов, формируемый инвентаризацией повторяющихся операций. Чтобы перечень не ограничивался тем, что владельцы процессов предъявляют сами, поиск ведётся систематически по основным функциональным областям предприятия (таблица 2); каждая область даёт типовые классы повторяющихся процессов, а документированные примеры из реестра рассмотренных кейсов (сопроводительные материалы; часть из них вошла в выборку раздела 5) показывают, что кандидаты существуют в каждой из них. Перечень двухуровневый: семь областей — сквозные и присутствуют на любом предприятии; отраслевое ядро заменяется по типу компании (для производственной — первые две строки, для торговой — торговые операции).

Таблица 2. Направления инвентаризации и документированные примеры

Направление	Типовые повторяющиеся процессы	Документированный пример
Основное производство	управление режимами агрегатов, дозирование компонентов, расчёт технологических параметров	Северсталь (температура прокатки), ММК (дозирование ферросплавов)
Обслуживание и ремонты	диагностика оборудования по телеметрии, планирование ремонтов	«Умный локомотив» (РЖД/ЛокоТех)
Торговые операции	прогноз спроса, контроль наличия на полке, приём заказов на потоке, складской фулфилмент	Ocado (частичный успех); Walmart и McDonald's (уроки откатов)
Логистика и снабжение	маршрутизация, планирование запасов, комплектация заказов	UPS ORION; DHL (комплектация)
Клиентский сервис и сбыт	обработка типовых обращений, суфлёр оператора	ассистент поддержки Fortune 500; Klarna (урок отката)
Финансы и учёт	сверки, контроль платежей, выявление мошенничества	Danske Bank (антифрод); сверка материального баланса

Направление	Типовые повторяющиеся процессы	Документированный пример
		(раздел 5.1)
Документооборот и договорная работа	извлечение реквизитов, проверка комплектности документов	JPMorgan (анализ кредитных договоров)
Инженерная инфраструктура	оптимизация энергопотребления, режимы инженерных систем	Google DeepMind (охлаждение ЦОД)
Управление и отчётность	консолидация данных, подготовка регулярной отчётности	перенос регулярной отчётности (раздел 5.1)
Персонал	скрининг резюме, типовые кадровые справки	Amazon и iTutorGroup (уроки: недоказуемость отсутствия смещения; судебное урегулирование)

Эта оговорка стоит на шаге инвентаризации не случайно: направления с дорогой ошибкой регулярно дают привлекательных по расчётному эффекту кандидатов (уроки в таблице 2), и без неё воронка отбора втянула бы их на общих основаниях. В направлениях, где ошибка дорога — решения о людях, юридически значимые и клинические, — атрибут «цена ошибки» оценивается по нижней границе, а ИИ запускается только в формате советчика: система предлагает, решение принимает человек; расширение автономии возможно лишь по правилу заработанной автономии (раздел 3.5).

Целиком модель применяется как последовательность из пяти шагов. (1) Инвентаризация по направлениям таблицы 2 даёт перечень кандидатов. (2) По каждому кандидату владелец процесса заполняет анкету четырнадцати атрибутов (приложение); значение 0 по любому из четырёх стоп-условий снимает кандидата. (3) Для прошедших вычисляется индекс осуществимости F и применяются зоны решения (раздел 3.3). (4) Кандидаты зоны $F \geq 70$ ранжируются индексом приоритета Π (раздел 3.4); в работу берётся один-два верхних. (5) Для выбранного фиксируются базовый замер и порог приёмки, пилот запускается на ступени автономии, соответствующей цене ошибки, и расширяет её по мере подтверждённого качества (раздел 3.5). Сквозной прогон всей последовательности на примере предприятия — раздел 5.1; но прежде чем показывать модель в работе, зафиксируем её отношение к ближайшим аналогам (раздел 4).

4. Сопоставление с рубрикой SML

Ближайший аналог модели — рубрика SML [Brynjolfsson, Mitchell, 2017; Brynjolfsson et al., 2018], и самое предсказуемое возражение к предлагаемой модели — «это переупакованная рубрика SML». Чтобы новизна была проверяемой, а не декларируемой, проводим попунктное сопоставление: четырнадцать атрибутов против восьми критериев и 21-пунктной рубрики (по [Septiandri et al., 2024]). Сопоставление разделяет атрибуты на три класса: что модель наследует, что добавляет и что осознанно отбрасывает.

Класс А — операционализация (семь атрибутов). Атрибуты 1–6 и 9 переводят критерии SML из крауд-оценки абстрактных описаний задач O*NET в инженерную проверку конкретного процесса. Смена оценивающего субъекта принципиальна: вместо краудворкера, оценивающего чужую абстракцию (что даёт узкое распределение оценок и слабую предсказательную силу [Acemoglu et al., 2022]), — владелец процесса, оценивающий собственную операцию по наблюдаемым фактам.

Класс Б — атрибуты без аналогов. Точка встраивания (8), владелец-спонсор (10), доступность экспертов (11) и определённая техническая реализация (13) не имеют соответствий ни в SML, ни в мерах подверженности [Felten et al., 2021; Webb, 2020; Tolan et al., 2021]; атрибут 14 (стабильность правил) прямо операционализирует критерий SML о неизменности феномена во времени, перенося его с абстрактной задачи на конкретный процесс. Два атрибута занимают промежуточное положение: контрольные случаи (7) частично соотносятся с критерием обратной связи SML, но здесь это требование наличия проверочного набора у конкретного процесса и бинарное стоп-условие; изменяемость регламента (12) соотносится с критерием стабильности феномена, но касается права организации менять регламент участка, а не природы задачи, и также переведена в стоп-условие. Атрибут 9, отнесённый к классу А, дополнительно играет отсутствующую в SML роль слагаемого экономической формулы Е. Именно организационно-интеграционных измерений недостаёт рубрике по признанию её авторов [Brynjolfsson et al., 2018], и их отсутствие объясняет слабую связь мер пригодности с фактическим внедрением (см. раздел 1.1) [Meindl et al., 2021].

Класс В — осознанно опущенное. Пункты рубрики, не относящиеся к корпоративным информационным процессам, в модель не вошли: физическая ловкость и мобильность, скорость реакции «менее секунды» (артефакт supervised-парадигмы 2017 г.), восприятие результата «как от человека»; требование объяснимости учтено косвенно — через атрибут цены ошибки и правило заработанной автономии.

Помимо состава критериев модель отличается механизмом принятия решения: некомпенсаторные стоп-условия вместо усреднения; расчётная экономическая формула вместо экспертного балла; пороги принятия решения вместо непрерывного индекса без правила решения; связь отбора с условиями реализации вместо статичной оценки *ex ante*. Модель, таким образом, не дублирует рубрики пригодности: её назначение — преодоление задокументированного разрыва «техническая пригодность $\neq$ фактическое внедрение».

Отличия от ближайших процедурных моделей также проверяемы по пунктно. Метод [Hofmann et al., 2020] генерирует кандидатов в пять шагов, но критерии оценки не операционализирует — организация вырабатывает их сама на этапе сравнения и приоритизации; предлагаемая модель даёт готовую шкалу с якорями. Подход [Bodendorf, 2025] оценивает кейсы по осям ценности (стратегическая и финансовая — экспертным баллом) и осуществимости (сложность данных, модели, экспертиза, интеграция, риск) с компенсаторным агрегированием, в расчёте на корпоративный портфель с выделенными ролями; предлагаемая модель заменяет балл ценности расчётной формулой E , вводит некомпенсаторные стоп-условия и рассчитана на владельца процесса без выделенных ролей. Обзор [Kakuschke et al., 2025] систематизирует критерии по осям желательности, осуществимости и жизнеспособности как основу для будущих инструментов, констатируя отсутствие самих инструментов. Блоки предлагаемой модели соотносятся с этими осями (технический блок и блок данных — осуществимость, экономика — жизнеспособность, организационный блок — желательность и готовность); сама модель относится к тому классу прикладных инструментов, дефицит которых обзор фиксирует. Сопоставление, однако, устанавливает лишь отличимость модели от аналогов, но не её работоспособность — последняя проверяется отдельно, демонстрацией и ретроспективной оценкой в разделе 5.

5. Демонстрация и оценка

5.1. Демонстрация

Демонстрация и оценка отвечают на разные вопросы, поэтому идут последовательно. Демонстрация показывает, что процедуру из пяти шагов (раздел 3.6) можно выполнить на реальном материале от начала до конца — от инвентаризации до зафиксированного порога приёмки — и на каждом шаге получить осмысленный результат; заодно она покрывает экономический блок (эффект E и приоритет Π), который ретроспективным кодированием не проверить. Оценка (5.2–5.3) отвечает на другой вопрос: различает ли модель успехи и неудачи на независимых кейсах. Работа модели демонстрируется на иллюстративном сценарии, построенном по материалам реального проекта на производственном предприятии; параметры обезличены и округлены. Перечень сформирован инвентаризацией повторяющихся операций по направлениям таблицы 2 совместно владельцами процессов и ведущим внедрения: 14 кандидатов, стоп-условия отсекали девять. Первую позицию ранжирования заняла ежедневная сверка материального баланса: выполняется 250 раз в год (N), занимает 6 человеко-часов (t), вход полностью оцифрован, владелец — начальник планового отдела; реализация отработана, правила сверки стабильны — анкета из четырнадцати атрибутов даёт индекс осуществимости $F = 89$. Годовой эффект по формуле раздела 3.4: при ставке 800 руб./час (c) и достижимой доле улучшения 0,8 (k — обоснована полнотой правил сверки, подтверждённой экспертами) $E \approx 0,96$ млн руб. в год; потери от невыявленных расхождений (слагаемое $L \times k'$) до пилота в

деньгах не оценивались и в расчёт не вошли. Стоимость внедрения — класс S, недели силами своей команды, то есть $C = 1$; итого приоритет $P = 0,96 \times 0,89 / 1 \approx 0,85$. Ближайший кандидат — перенос регулярной отчётности: эффект сопоставим ($E \approx 0,9$ млн руб.), но осуществимость ниже ($F = 75$), а внедрение дороже (класс M, $C = 3$), поэтому $P \approx 0,22$ — почти вчетверо меньше; разрыв сохраняется и при другом кодировании стоимости ($S/M/L = 1/2/4$ даёт 0,85 против 0,34). Полный разбор заполнения анкеты для этого сценария — в приложении. Базовая линия зафиксирована до старта (6 человеко-часов в день; до четырёх расхождений в неделю уходили в себестоимость незамеченными), порог приёмки согласован заранее (обнаружение $\geq 95\%$ расхождений, подтверждённых контролёром). Контрастный кандидат — согласование договоров — при высокой ожидаемой ценности отсечён стоп-условием 12 (регламент задан внешними требованиями): демонстрация некомпенсаторности, отличающей модель от балльных методов, где такой кандидат прошёл бы за счёт ценности.

5.2. Дизайн ретроспективной оценки

Оценка построена в два этапа — так рамка FEDS [Venable et al., 2016] предписывает проверять методы, ошибка которых затрагивает людей и деньги: сначала ранняя проверка на доступном материале (она и выполнена в статье), затем полноценная проверка на реальных внедрениях по мере их накопления; малое число кейсов на первом этапе — норма этой рамки, а не дефект. Для ранней проверки выполнено ретроспективное кодирование независимой выборки из 14 опубликованных внедрений ИИ с задокументированными исходами: 7 неудач/откатов и 7 успехов (состав — в таблице 3). Критерий включения — наличие первичных документов, рецензируемых публикаций или публичных описаний с измеренными величинами (для двух российских кейсов — публикации компаний в отраслевой прессе), позволяющих кодировать состояние процесса до исхода. Среди них — аудит University of Texas System по проекту MD Anderson, биржевые отчёты Zillow Group, решение трибунала по делу Air Canada, независимая клиническая валидация Epic Sepsis Model [Wong et al., 2021], рецензируемые описания успехов [Holland et al., 2017; Brynjolfsson et al., 2025; Gomez-Urbe, Hunt, 2015].

Главная угроза такого дизайна — ретроспективное искажение (hindsight bias): кодировщик, знающий исход, невольно подгоняет оценки под него, и согласованность модели с исходами становится артефактом знания исходов, а не свойством модели. Против этого протокол предусматривает: (1) правило отсечки — атрибуты кодируются только по фактам о состоянии процесса до исхода; источники, опубликованные после исхода, используются исключительно для извлечения фактов о состоянии «до», кодирование причинных интерпретаций из них запрещено; (2) обязательность доказательства — каждый код сопровождается фактом с источником, при отсутствии факта атрибут не кодируется; (3) двойное независимое кодирование со сверкой и фиксацией согласованности; (4) если атрибут по публичным данным закодировать не удалось, индекс

считается по закодированным — в предположении, что пропуски не связаны со значениями атрибутов.

Оценка выполнена по протоколу двойного независимого кодирования: каждый кейс кодировался дважды и независимо, с разными установками (консервативной и буквалистской), слепым к гипотезам валидации способом; расхождения разрешались по силе доказательства. Автор в кодировании не участвовал. Включённые кейсы в ряде случаев крупнее целевого объекта модели — вплоть до целых бизнес-направлений (Zillow Offers, Cruise); это компромисс в пользу документированности исходов, ограничивающий перенос выводов на рутинные корпоративные процессы (раздел 5.4). Выборка намеренно охватывает обе технологические эпохи — классическое машинное обучение (UPS, Danske Bank, Netflix) и большие языковые модели (ассистенты Fortune 500 и Klarna): атрибуты сформулированы как свойства процесса и организации, а не возможностей конкретного поколения моделей, и кодируются одинаково в обеих эпохах.

5.3. Результаты

Столбец первичного источника включён в таблицу результатов не для справки: доказательная сила каждого кода определяется типом документа — исходы неудач зафиксированы независимыми инстанциями (аудит, суд, биржевая отчётность, трибунал), а не пересказами, и каждый код проверяем по реестру источников.

Таблица 3. Результаты ретроспективного кодирования (F — индекс осуществимости по закодированным атрибутам; стопы — сработавшие стоп-условия)

Кейс	Исход	Первичный источник	F	Сработавшие стопы
MD Anderson + IBM Watson	остановлен после \$62 млн затрат	аудит University of Texas System, 2017	38	7, 8
Zillow Offers	закрыт, списания >\$0,5 млрд	отчётность SEC (8-K, 10-K), 2021	75	— (атрибут 7 не кодируем)
Klarna, ассистент поддержки на большой языковой модели	откат, возврат людей	пресс-релиз компании, 2024; Bloomberg, 2025	73	7
Air Canada, чат-бот	проигранное разбирательство	решение трибунала 2024 BCCRT 149	59	7, 8, 10
Epic Sepsis Model	опровергнута независимой валидацией	[Wong et al., 2021]	71	7 (атрибут 10 не кодируем)

Кейс	Исход	Первичный источник	F	Сработавшие стопы
GM Cruise, роботакси	закрыт, инвестировано >\$10 млрд	решения DMV Калифорнии, 2023; заявление GM, 2024	46	12
Commonwealth Bank, голосовой бот	отмена решения, признание ошибки	материалы Fair Work Commission, 2025	55	7
UPS ORION	успех, \$300–400 млн/год	[Holland et al., 2017]	88	—
Fortune 500, ассистент поддержки	успех, +15% производительности	[Brynjolfsson et al., 2025]	81	—
Danske Bank, антифрод	успех	публикации Teradata и Forbes, 2017	89	—
Google DeepMind, охлаждение ЦОД	успех, –40% энергии на охлаждение	публикация Google/DeepMind, 2016	81	—
Netflix, рекомендации	успех	[Gomez-Uribe, Hunt, 2015]	92	—
Северсталь, стан 2000	успех, +24 тыс. т за 3 мес.	отраслевая пресса (ComNews), 2023	85	—
ММК, дозирование ферросплавов	успех, эффект 2,7–4%	отраслевая пресса (JetInfo), 2020	77	—

Полный реестр источников кодирования (адреса, даты публикации, маркеры «до/после исхода») приведён в сопроводительных материалах.

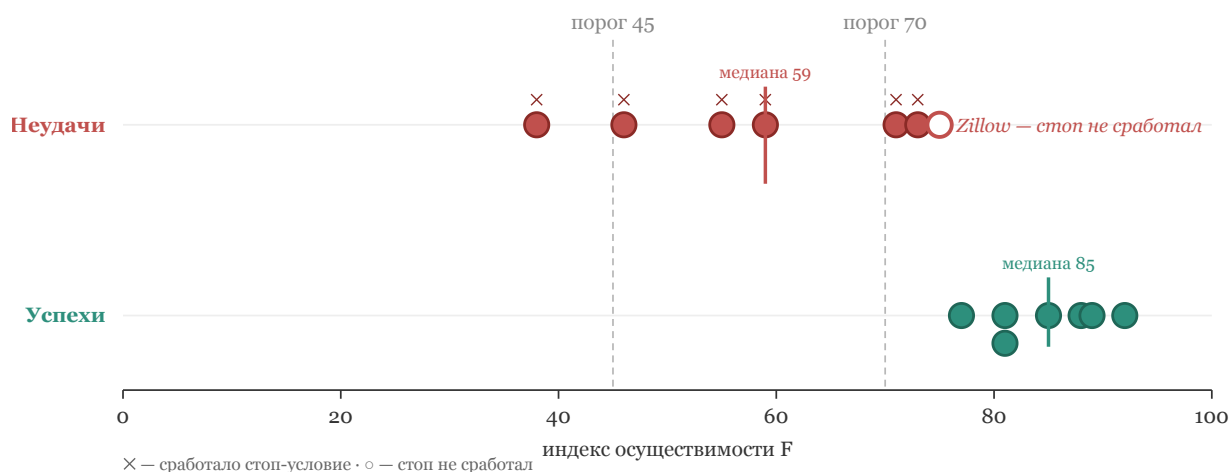


Рисунок 1. Распределение индекса осуществимости F по группам. Успехи сосредоточены в зоне выше порога 70; неудачи рассеяны ниже, шесть из семи со сработавшим стоп-условием. Единственная неудача без сработавшего стопа — Zillow (дрейф данных после запуска) — граничный случай.

Модель различает группы. У шести из семи неудач сработало хотя бы одно стоп-условие; у успехов — ни одного из 28 проверенных (все стоп-атрибуты успехов закодированы; об асимметрии источников успехов — раздел 5.4). Самое частое сработавшее стоп-условие — «контрольные случаи для замера» (пять кейсов из шести): организации запускали и масштабировали системы без определённых заранее критериев приёмки — от двенадцати продлений контракта у MD Anderson до тиражирования Epic Sepsis Model в сотни больниц до независимой валидации. Медианный индекс осуществимости группы неудач — 59 против 85 у группы успехов; все семь успехов имеют $F \geq 70$. Совокупное отсекающее правило «сработавшее стоп-условие или $F < 45 \rightarrow$ не брать» отсекает шесть из семи неудач и пропускает все семь успехов. Случайно такое разделение получить почти невозможно: точный критерий Фишера — стандартная проверка неслучайности для малых таблиц — оценивает вероятность случайного совпадения в 0,2% ($p = 0,0023$; двусторонний вариант — 0,0047). Значение приводится справочно: выборка не случайна, поэтому переносить эти доли на все внедрения ИИ нельзя. Отметим, что ретроспективно проверена именно отсекающая часть правила: предписываемый практику фильтр $F \geq 70$ в чистом виде пропустил бы и Klarna (73), и Zillow (81) — различие групп обеспечивают преимущественно стоп-условия, вклад порогов индекса вторичен. При этом классификация выборки устойчива к положению порога отсечения: его сдвиг в любой точке диапазона 42–73 результата не меняет (ближайшие значения F — 41 у MD Anderson и 73 у ММК), чувствительность появляется лишь у порогов выше минимального F успехов.

Два кейса содержательно важнее сводных долей. Кейс Klarna ($F = 73$ — выше порога) отсекается только стоп-условием 7: метрика качества обслуживания не входила в контур решения, оптимизировалась стоимость контакта; компенсаторная балльная модель

пропустила бы этого кандидата — некомпенсаторность стоп-условий содержательно необходима. Показательна пара с близким индексом (Klarna 73, ММК 77): успех ММК и откат Klarna различает именно это стоп-условие — ММК задал проектную метрику эффекта (3%) до опытно-промышленной эксплуатации и сверил её с фактической (2,7–4% по сортаменту), тогда как у Klarna контрольного замера качества не было. При равной суммарной осуществимости судьбу внедрения решило наличие метрики. Единственное расхождение — Zillow Offers ($F = 75$; три кодируемых стоп-условия пройдены, атрибут 7 по публичным источникам не кодируем): у процесса были владелец и полная оцифрованность, а провал наступил из-за дрейфа данных после запуска и накопления убытка до момента обнаружения. Показательно, что атрибут стабильности правил (14) оказался нулевым только у Zillow среди всех четырнадцати кейсов — он изолирует именно тот механизм (нестационарность рынка жилья), которым кейс и погиб, снижая индекс с 81 до 75, хотя и не переводя его за порог. Расхождение было предсказано при формировании выборки до кодирования и очерчивает границу применимости: модель оценивает осуществимость на входе, но не заменяет контроль качества после запуска — за это в модели отвечают правила базовой линии и заработанной автономии, а не индекс осуществимости.

Два кодировщика совпали в 87% оценок (125 из 144 пар, где оба вынесли суждение; в остальных 24 хотя бы один воздержался за отсутствием данных). С поправкой на случайные совпадения — её даёт каппа Коэна, стандартная мера согласия двух оценщиков — согласованность равна 0,76 (взвешенная — 0,80), что по принятой шкале считается существенной. Расхождения разрешались по правилу силы доказательства с письменной фиксацией. Атрибут 4 (доля исключений) оказался систематически не кодируемым по публичным источникам — закодировать его удалось лишь в одном кейсе из четырнадцати. Это ожидаемое свойство: величина наблюдаема для владельца процесса, но не для внешнего наблюдателя, что подтверждает выбор владельца процесса как оценивающего субъекта.

Результат интерпретируется строго как ретроспективная согласованность: модель, применённая к фактам, известным до исхода, не противоречит задокументированным исходам, а её стоп-условия соответствуют реально сработавшим механизмам неудач. Это необходимое, но не достаточное условие валидности; утверждения о предсказательной точности из данного дизайна не следуют.

5.4. Ограничения оценки

Ограничения оценки состоят в следующем. Ретроспективное искажение смягчено протоколом на уровне кодирования, но не на уровне формирования выборки: кейсы отбирал исследователь, знавший исходы; смягчения — фиксированный критерий включения и публикация полного реестра рассмотренных кандидатов (57 кейсов) в сопроводительных материалах, так что отвергнутые кандидаты доступны проверке. Есть и

круг в самой конструкции: атрибуты выведены из литературы о причинах неудач — и проверяются на кейсах, чьи причины та же литература и задокументировала. Поэтому результат показывает прежде всего, что известные диагнозы удалось перевести в проверяемые признаки конкретного процесса, — а не независимое подтверждение состава атрибутов. Выборка не случайна: в публичное поле попадают резонансные неудачи и успехи с публикациями самих компаний, поэтому судить по ней о том, как часто внедрения проваливаются вообще, нельзя. Источники асимметричны: неудачи документированы независимыми аудитами и судами, успехи — чаще самими компаниями, что систематически завышает оценки атрибутов успехов; для каждого кода фиксировался тип доказательства, а тип источника (первичный документ, пресса, вендорский самоотчёт) — в реестре источников. При $n = 14$ статистические выводы неправомерны — сравнение групп носит описательный характер. Ретроспективная оценка покрыла индекс осуществимости и стоп-условия; экономический блок (Е, П) проверен только демонстрацией. Целевое свойство заполняемости за одну рабочую сессию в настоящей работе не оценивалось; его проверка (время заполнения, сходимость оценок двух независимых владельцев, устойчивость к заинтересованному завышению оценок владельцем-спонсором) — обязательная часть проспективной валидации. Наконец, ретроспективная согласованность не равна прогностической силе: её подтверждение требует проспективного применения модели к пилотам, отобраным до старта.

6. Обсуждение

Зафиксированные ограничения очерчивают границы, в которых статья вправе формулировать выводы; внутри этих границ вклад исследования: (1) первый в рассмотренной литературе компактный инструмент отбора уровня отдельного процесса, соединяющий техническую осуществимость, расчётную экономику эффекта, организационную готовность и условия реализации — существующие альтернативы покрывают эти измерения порознь; (2) формальное введение в модель отбора ИИ-инициатив некомпенсаторных стоп-условий — механизма, отсутствующего и в рубриках пригодности, и в компенсаторных MCDM-фреймворках и при этом соответствующего реальным механизмам неудач; (3) перевод критериев SML из крауд-оценки абстрактных задач в анкету, которую заполняет владелец конкретного процесса. Для русскоязычного поля модель является первой критериальной моделью отбора ИИ-инициатив, стыкующейся с линией исследований зрелости [Лобокский и др., 2026]: их предмет — как внедрять после решения, предмет модели — что и в каком порядке отбирать до него.

Ограничения модели — продолжение её проектных решений. Равные веса атрибутов и пороги 70/45 калиброваны экспертно; их эмпирическая калибровка требует накопления когорты внедрений; до неё пороги 70/45 следует трактовать как ориентиры, а не жёсткие границы, опираясь в первую очередь на стоп-условия — именно их различающую способность, а не положение порогов, подтвердила оценка (раздел 5.3).

Атрибуты технического блока могут коррелировать между собой, а цена ошибки участвует в модели дважды — как фильтр пригодности в индексе и как слагаемое потерь в эффекте; обе особенности требуют проверки на накопленных данных. Рост возможностей моделей со временем смещает оценки атрибутов; большинство атрибутов сформулировано технологически нейтрально, однако атрибут 3 перечисляет классы работ, посильные текущему поколению технологий, и подлежит пересмотру при их развитии; периодическая рекалибровка остаётся штатной процедурой сопровождения, а правило заработанной автономии — встроенным механизмом адаптации.

Заключение

В статье предложена формализованная модель отбора процессов предприятия для внедрения ИИ: четырнадцать атрибутов со шкалой 0/1/2, четыре некомпенсаторных стоп-условия, индекс осуществимости с порогами принятия решения и индекс приоритета с расчётной формулой эффекта. Ретроспективное кодирование четырнадцати опубликованных кейсов показало согласованность индекса и стоп-условий с задокументированными исходами: сработавшие стоп-условия у шести из семи неудач при нуле у успехов и разрыв медианных индексов 59 против 85. Направления дальнейших исследований: проспективная валидация — применение модели к пилотам до старта с последующим сопоставлением прогноза и исхода; эмпирическая калибровка весов и порогов по когорте внедрений; независимое слепое перекодирование атрибутов реализуемости (13–14); эмпирическая проверка правила быстрых побед; расширение модели на агентные системы, где единицей отбора становится не отдельный процесс, а связка процессов с передачей контекста.

Литература

1. Белова Е.Ю., Бадьина А.В. (2025) Измерение эффекта от внедрения ИИ на предприятиях // Прогрессивная экономика. № 3. С. 42–59.
2. Гилева Т.А. (2019) Цифровая зрелость предприятия: методы оценки и управления // Вестник УГНТУ. Наука, образование, экономика. Серия: Экономика. № 1 (27). С. 38–52.
3. Грекул В.И., Исаев Е.А., Коровкина Н.Л., Лисиенкова Т.С. (2019) Разработка подхода к ранжированию инновационных ИТ-проектов // Бизнес-информатика. Т. 13. № 2. С. 43–58.
4. Долганова О.И., Деева Е.А. (2019) Готовность компании к цифровым преобразованиям: проблемы и диагностика // Бизнес-информатика. Т. 13. № 2. С. 59–72.
5. Комаров Н.М., Голубев С.С., Пашенко Д.С., Щербаков А.Г. (2024) Инструменты искусственного интеллекта в программах цифровой трансформации промышленных предприятий // Мир новой экономики. Т. 18. № 3. С. 6–16. DOI: 10.26794/2220-6469-2024-18-3-6-16.
6. Краковская И.Н., Корокошко Ю.В., Слушкина Ю.Ю. (2024) Цифровая зрелость промышленных предприятий: опыт оценки // Вестник Санкт-Петербургского

- университета. Экономика. Т. 40. № 3. С. 433–459. DOI: 10.21638/spbu05.2024.305.
7. Лобочкий М.Ю., Комаров М.М., Зараменских Е.П. (2026) Maturity model and success factors for the implementation of AI agents in Russian companies // *Business Informatics*. Vol. 20. No. 2. P. 7–21.
 8. Марон А.И., Файнбург Г.Д. (2025) Определение последовательности реализации проектов программы повышения эффективности бизнес-процессов // *Бизнес-информатика*. Т. 19. № 1. С. 22–33. DOI: 10.17323/2587-814X.2025.1.22.33.
 9. Acemoglu D., Autor D., Hazell J., Restrepo P. (2022) Artificial Intelligence and Jobs: Evidence from Online Vacancies // *Journal of Labor Economics*. Vol. 40. No. S1. P. S293–S340.
 10. Agrawal A., Gans J., Goldfarb A. (2018) *Prediction Machines: The Simple Economics of Artificial Intelligence*. Boston: Harvard Business Review Press.
 11. Ali W., Khan A.Z. (2025) Factors influencing readiness for artificial intelligence: a systematic literature review // *Data Science and Management*. Vol. 8. No. 2. P. 196–206. DOI: 10.1016/j.dsm.2024.09.005.
 12. Archer N.P., Ghasemzadeh F. (1999) An integrated framework for project portfolio selection // *International Journal of Project Management*. Vol. 17. No. 4. P. 207–216.
 13. Arntz M., Gregory T., Zierahn U. (2016) *The Risk of Automation for Jobs in OECD Countries*. OECD Social, Employment and Migration Working Papers No. 189. Paris: OECD Publishing.
 14. Bodendorf F. (2025) A data-driven use case planning and assessment approach for AI portfolio management // *Electronic Markets*. Vol. 35. Article 22.
 15. Brynjolfsson E., Mitchell T. (2017) What can machine learning do? Workforce implications // *Science*. Vol. 358. No. 6370. P. 1530–1534.
 16. Brynjolfsson E., Mitchell T., Rock D. (2018) What Can Machines Learn, and What Does It Mean for Occupations and the Economy? // *AEA Papers and Proceedings*. Vol. 108. P. 43–47.
 17. Brynjolfsson E., Li D., Raymond L. (2025) Generative AI at Work // *The Quarterly Journal of Economics*. Vol. 140. No. 2. P. 889–942.
 18. Davenport T.H., Ronanki R. (2018) Artificial Intelligence for the Real World // *Harvard Business Review*. Vol. 96. No. 1. P. 108–116.
 19. Eloundou T., Manning S., Mishkin P., Rock D. (2024) GPTs are GPTs: Labor market impact potential of LLMs // *Science*. Vol. 384. No. 6702. P. 1306–1308.
 20. Felten E.W., Raj M., Seamans R. (2021) Occupational, Industry, and Geographic Exposure to Artificial Intelligence // *Strategic Management Journal*. Vol. 42. No. 12. P. 2195–2217.
 21. Frank M.R., Autor D., Bessen J.E. et al. (2019) Toward understanding the impact of artificial intelligence on labor // *PNAS*. Vol. 116. No. 14. P. 6531–6539.
 22. Frey C.B., Osborne M.A. (2017) The future of employment: How susceptible are jobs to computerisation? // *Technological Forecasting and Social Change*. Vol. 114. P. 254–280.
 23. Gartner (2024) Gartner Predicts 30% of Generative AI Projects Will Be Abandoned After Proof of Concept by End of 2025. Press release, 29 July 2024. Stamford: Gartner.
 24. Georgieff A., Milanez A. (2021) What happened to jobs at high risk of automation? OECD Social, Employment and Migration Working Papers No. 255. Paris: OECD Publishing.

25. Gomez-Uribe C.A., Hunt N. (2015) The Netflix Recommender System: Algorithms, Business Value, and Innovation // *ACM Transactions on Management Information Systems*. Vol. 6. No. 4. Article 13.
26. Gregor S., Hevner A.R. (2013) Positioning and Presenting Design Science Research for Maximum Impact // *MIS Quarterly*. Vol. 37. No. 2. P. 337–355.
27. Hevner A.R. (2007) A Three Cycle View of Design Science Research // *Scandinavian Journal of Information Systems*. Vol. 19. No. 2. Article 4.
28. Hevner A.R., March S.T., Park J., Ram S. (2004) Design Science in Information Systems Research // *MIS Quarterly*. Vol. 28. No. 1. P. 75–105.
29. Hofmann P., Jöhnk J., Protschky D., Urbach N. (2020) Developing Purposeful AI Use Cases — A Structured Method and Its Application in Project Management // *Proceedings of the 15th International Conference on Wirtschaftsinformatik*. Potsdam.
30. Holland C., Levis J., Nuggehalli R., Santilli B., Winters J. (2017) UPS Optimizes Delivery Routes // *Interfaces*. Vol. 47. No. 1. P. 8–23.
31. Holmström J. (2022) From AI to digital transformation: The AI readiness framework // *Business Horizons*. Vol. 65. No. 3. P. 329–339.
32. Jöhnk J., Weißert M., Wyrski K. (2021) Ready or Not, AI Comes — An Interview Study of Organizational AI Readiness Factors // *Business & Information Systems Engineering*. Vol. 63. No. 1. P. 5–20.
33. Kakuschke N., Legner C., Jung R. (2025) Creating Value from Data the Right Way — A Framework for Assessing Data Value Creation Use Cases // *Proceedings of the 58th Hawaii International Conference on System Sciences*.
34. March S.T., Smith G.F. (1995) Design and natural science research on information technology // *Decision Support Systems*. Vol. 15. No. 4. P. 251–266.
35. Meindl B., Frank M.R., Mendonça J. (2021) Exposure of occupations to technologies of the fourth industrial revolution. arXiv:2110.13317.
36. MIT NANDA (2025) The GenAI Divide: State of AI in Business 2025. Preliminary report. Cambridge, MA: MIT Project NANDA.
37. Nedelkoska L., Quintini G. (2018) Automation, skills use and training. OECD Social, Employment and Migration Working Papers No. 202. Paris: OECD Publishing.
38. Paleyes A., Urma R.-G., Lawrence N.D. (2022) Challenges in Deploying Machine Learning: A Survey of Case Studies // *ACM Computing Surveys*. Vol. 55. No. 6. Article 114.
39. Peffers K., Tuunanen T., Rothenberger M.A., Chatterjee S. (2007) A Design Science Research Methodology for Information Systems Research // *Journal of Management Information Systems*. Vol. 24. No. 3. P. 45–77.
40. Ryseff J., De Bruhl B.F., Newberry S.J. (2024) The Root Causes of Failure for Artificial Intelligence Projects and How They Can Succeed. Santa Monica: RAND Corporation.
41. Sculley D., Holt G., Golovin D. et al. (2015) Hidden Technical Debt in Machine Learning Systems // *Advances in Neural Information Processing Systems 28 (NIPS 2015)*.

42. Septiandri A.A., Constantinides M., Quercia D. (2024) The potential impact of AI innovations on US occupations // PNAS Nexus. Vol. 3. No. 9.
43. Tolan S., Pesole A., Martínez-Plumed F. et al. (2021) Measuring the Occupational Impact of AI: Tasks, Cognitive Abilities and AI Benchmarks // Journal of Artificial Intelligence Research. Vol. 71. P. 191–236.
44. Venable J., Pries-Heje J., Baskerville R. (2016) FEDS: a Framework for Evaluation in Design Science Research // European Journal of Information Systems. Vol. 25. No. 1. P. 77–89.
45. Webb M. (2020) The Impact of Artificial Intelligence on the Labor Market. SSRN Working Paper No. 3482150.
46. Westenberger J., Schuler K., Schlegel D. (2022) Failure of AI projects: understanding the critical factors // Procedia Computer Science. Vol. 196. P. 69–76.
47. Wong A., Otles E., Donnelly J.P. et al. (2021) External Validation of a Widely Implemented Proprietary Sepsis Prediction Model in Hospitalized Patients // JAMA Internal Medicine. Vol. 181. No. 8. P. 1065–1070.

Приложение. Анкета отбора процесса и пример заполнения

Приложение — рабочая форма модели, применимая без обращения к основному тексту: таблица A1 используется как анкета на шаге 2 порядка применения (раздел 3.6), таблица A2 задаёт эталон заполнения — уровень конкретности фактов-обоснований, который ведущий внедрения вправе требовать от владельца процесса.

Таблица A1. Анкета: двенадцать атрибутов с якорями оценок

№	Атрибут	2	1	0
1	Однотипность выполнения	каждое выполнение похоже на предыдущее	несколько типовых вариантов	каждый раз по-разному
2	Формализуемость правил	есть письменная инструкция	правила в головах экспертов	правил нет
3	Класс работы	сопоставление / извлечение / классификация / расчёт / контроль по изображению / генерация текста	смешанная: часть — ручные действия	физические действия, переговоры, ответственные решения
4	Доля исключений («не по правилам»)	меньше 10%	10–30%	больше 30%

№	Атрибут	2	1	0
5	Оцифрованность входа	полностью в системах	частично; есть бумага/файлы	бумага, голос, «в головах»
6	История выполнения с правильными результатами	годы истории в системах	немного примеров	истории нет
7	Контрольные случаи для замера (стоп)	соберём 100+ случаев с известным ответом	несколько десятков	собрать не можем
8	Точка встраивания (стоп)	есть система / API / файловый обмен	потребуется доработка	встроиться некуда
9	Цена ошибки	низкая: ошибка легко ловится и исправляется	средняя: стоит денег, но обратима	высокая: необратимый ущерб или безопасность
10	Владелец-спонсор (стоп)	активно заинтересован в результате	нейтрален	владельца нет
11	Эксперты на верификацию	выделены регулярные часы	от случая к случаю	недоступны
12	Регламент допускает изменение исполнения (стоп)	да	через согласование	нет, участок неприкасаем
13	Определённость технической реализации (разведка)	отработанные решения, аналоги документированы	нужна адаптация или дообучение	способ решения неясен
14	Стабильность правил процесса	стабильны годами	предсказуемый цикл изменений	меняются быстрее выхода на порог

Правила обработки: атрибут 3 оценивается по составу работы (2 — процесс целиком из перечисленных информационных операций; 1 — они перемешаны с ручными действиями; 0 — ядро процесса — физические действия, переговоры или ответственные решения); атрибут 4:

исключение — выполнение, не укладывающееся в типовые правила и требующее нестандартного разбора человеком, доля — от общего числа выполнений; атрибут 6: история засчитывается, только если для прошлых выполнений известен подтверждённый правильный итог (пары «вход — верный выход»), журнал без итогов не годится; атрибут 11: считается не наличие экспертов, а выделенные им регулярные часы на верификацию знаний, контрольные случаи и разбор ошибок; значение 0 по атрибуту 7, 8, 10 или 12 — стоп, кандидат снимается с рассмотрения; $F = (\sum a_i / 28) \times 100$; зоны решения — $F \geq 70$, 45–70, $F < 45$ (раздел 3.3). Анкету заполняет владелец процесса; каждая оценка сопровождается фактом-обоснованием — оценки без фактов не принимаются ведущим внедрения (частичная защита от заинтересованного завышения).

Таблица А2. Пример заполнения: ежедневная сверка материального баланса (сценарий раздела 5.1; по материалам реального проекта, обезличено)

№	Атрибут	Оценка	Факт-обоснование
1	Однотипность выполнения	2	сверка выполняется ежедневно по одной схеме
2	Формализуемость правил	2	методика сверки закреплена в регламенте планового отдела
3	Класс работы	2	сопоставление учётных данных и расчёт расхождений
4	Доля исключений	1	10–30% позиций требуют ручного разбора причин расхождения
5	Оцифрованность входа	2	все данные ведутся в учётной системе
6	История выполнения	2	журналы сверок за годы работы доступны в системе
7	Контрольные случаи (стоп)	2	более ста прошлых сверок с подтверждённым результатом
8	Точка встраивания (стоп)	2	учётная система с программным доступом к данным

№	Атрибут	Оценка	Факт-обоснование
9	Цена ошибки	1	пропущенное расхождение искажает себестоимость, но исправимо при обнаружении
10	Владелец-спонсор (стоп)	2	начальник планового отдела активно заинтересован в результате
11	Эксперты на верификацию	1	экономисты доступны, но регулярные часы не выделены
12	Регламент (стоп)	2	регламент сверки внутренний, его изменение в компетенции предприятия
13	Определённость реализации	2	сверка правилами и расчётом — отработанный класс, аналоги документированы
14	Стабильность правил	2	методика сверки и номенклатура стабильны годами

$\Sigma a_i = 25$; $F = (25 / 28) \times 100 \approx 89$. Стоп-условия пройдены; кандидат проходит в ранжирование (расчёт Е и П — в разделе 5.1).