

♦ ПОД КАПОТОМ

175 страниц, которых никто не прочёл

Разбор двух редакций whitepaper Сбера «AI-Disrupt PDLC»: краткой на 44 полосы и полной на 175. Чем глубже копаешь, тем меньше доказательств. В краткой двенадцать источников, в полной ноль.

Источник: оба PDF с aipdlc.ru · автор по подписи

Кирилл Меньшов, CIO Сбера

Проверка: шесть параллельных агентов, 400+ обращений к первоисточникам, семь кругов ревью

Двадцать одна правка самого разбора, журнал в конце.

Всё непроверенное помечено явно

Два документа с одним названием

12

источников на 44 полосы краткой редакции. Для СМИ, ЦИПРа, Иннопрома и совета директоров

0

источников на 175 полос полной. Ни URL, ни сноски, ни раздела литературы

Обычно бывает наоборот: краткая версия ссылается на полную, где всё обосновано. Здесь полная обосновывает меньше краткой.

◆ ГЛАВНОЕ

Проблема не в том, что документ плох. В том, что он **не сведён сам с собой**: центральная шкала прав агента описана дважды и описания не сопоставлены, шкал зрелости пять вместо заявленных трёх, одни и те же числа в разных главах взяты из разных источников, а сроки внедрения из главы 3 и главы 7 расходятся в 1,5–2,5 раза без объяснения.

Документ читается как текст, не прошедший финальную редакторскую сверку: ключевые шкалы не сопоставлены между собой, сроки не сведены в критический путь, а числа не привязаны к работам, из которых взяты.

При этом внутри полной версии есть настоящая инженерная работа: формула пересчёта прав агента, система метрик с формулами и порогами, честный разбор конкурирующей методологии AWS. Ей посвящена **первая** часть разбора, до всех обвинений.

Что именно разбирается

	КРАТКАЯ	ПОЛНАЯ
Объём	44 полосы	175 полос
Где представлена	ЦИПР 19.05, Иннопром 07.07	aipdlc.ru
Источников	12	0
URL в тексте	—	0
Сносок	0	0
Растровых изображений	—	2 (обложечные)
GigaCode / GigaIDE	0 / 0	0 / 0
Автор по подписи	Кирилл Меньшов, CIO Сбера — колонтитул на каждой полосе	
Автор по метаданным	—	Альвианский Алексей, Word LTSC, 13.07.2026
Языки лендинга	русский · английский · 中文	

Про изображения. Проверено технически: pdfimages находит в 175-страничном PDF ровно два растровых объекта. Векторная графика есть, на 105 страницах из 175. Но выборочный рендер полос с максимальной плотностью объектов показывает разметку таблиц, а не диаграммы. Все 175 полос подряд никто не просматривал. Документ, который называет себя «новой архитектурой», описывает её прозой и таблицами.

Шесть частей

- I **Что в документе работает**
Формула пересчёта автономии, метрики с порогами, Evidence Bundle, честный разбор AWS AI-DLC

 - II **Чем глубже, тем меньше доказательств**
Ноль источников на 175 полос. 42% числовых утверждений без атрибуции вообще

 - III **Дорожная карта, по которой нельзя построить план**
Сроки расходятся в 1,5–2,5 раза между главами. Пятнадцать компонентов, одно число

 - IV **Документ, не сведённый сам с собой**
R0–R5 описана дважды, шкал L пять вместо трёх, одни числа — два источника

 - V **Цифры, которые не проверяются**
22 п.п. без источника, «1–3 п.п.» опровергается бенчмарками, CTOR не считается

 - VI **Правовая часть: два недействующих акта**
683-П отменено, приказ ФСТЭК 17 утратил силу, ст. 18.1 вместо ч. 5 ст. 18
-

В конце: пять вопросов для переговорной, сводка «брать / проверять / не верить», список непроверенного и журнал из двадцати одной правки самого разбора.



Что в документе работает

Раздел стоит первым намеренно. Всё, что идёт дальше, значит не «Сбер ничего не сделал», а «вот что сделано, и вот чем это испорчено».

Формула пересчёта автономии

Самая сильная инженерная конструкция документа. Уровень прав агента не задан статически, а пересчитывается на каждой операции.

ДОСЛОВНО

«**эффективный_R = номинальный_R – повышение_критичности – повышение_риска – повышение_сложности**, где повышение — число ступеней (0, 1 или 2), на которое снижается автономия из-за высокого значения фактора»

Факторов эскалации восемь: среда, ветка, критичность файла, обратимость действия, покрытие тестами, радиус воздействия, статус Evidence Bundle, длина горизонта.

◆ С ЧИСЛЕННЫМ ПРИМЕРОМ И РФ-СМЕЩЕНИЯМИ

«Задача, классифицированная как R3... высокая критичность (+2) и высокий профиль риска (+1). Тогда $\text{эффективный_R} = 3 - 2 - 1 = 0$ ». Фиксированные смещения под банковский контур: «изменение AML-конвейера — R2 → R0».

В проверенных публичных шкалах (SAE J3016, шкалы автономии агентов 2025–2026 годов) такой риск-адаптивной формулы не обнаружено. Это один из наиболее самостоятельных элементов документа.

Метрики с формулами и токеномика

Глава 6 остаётся единственным местом, где документ выполняет собственное обещание «для каждой метрики определены формула, целевые значения, частота и источник данных». Выполняет буквально.

МЕТРИКА	ФОРМУЛА	JUNIOR → SENIOR
Spec Reuse Rate	<code>reused_specs / total_specs</code>	≥10% → ≥50%
Hallucination Rate	<code>incorrect / total</code> на golden set	≤5% → ≤2%
Context Persistence	<code>consistent / total_long_sessions</code>	≥80% → ≥95%
Token Cost per Feature	<code>total_tokens / deployed_features</code>	×1 → ×0.4

♦ **COST CIRCUIT BREAKERS – ТРИ УРОВНЯ**

Бюджеты на задачу, на сессию и на горизонт: **50К / 200К / 500К** токенов; при превышении 80% подтверждение запрашивается автоматически.

Это единственное место во всём документе, где происхождение чисел названо честно: «калибровочные значения по внутренним наблюдениям Сбера, настраиваются под организацию».

Плюс раздел антипаттернов измерения, где закон Гудхарта назван по имени, а vanity metrics разобраны отдельно. Для корпоративного документа это редкость.

Evidence Bundle и честность к конкуренту

В **Evidence Bundle** входит восемь обязательных элементов: diff со ссылкой на PR, результаты четырёх классов проверок, отчёт о расходе токенов, append-only журнал сессии, аттестации четырёх шлюзов, структурированное обоснование для операций R2+, вклад в объём результата, указатель на измерение бизнес-исхода. Пороги заданы: регрессии не ниже 95%, покрытие критериев приёмки 100%.

♦ РАЗДЕЛ «AWS AI-DLC И AI-DISRUPT PDLC: СХОДСТВА И РАЗЛИЧИЯ»

AWS AI-DLC назван «наиболее проработанной промышленной методологией ИИ-нативной разработки». Перечислены три конкретных отличия: AI-DLC не выделяет Discovery как обязательный этап, не различает SDD и Agent Specification, не формализует контур валидации.

Это ровно то, чего не хватает всему остальному документу. Здесь оно есть.

РАЗЛИЧИЕ «КОНТЕКСТ ПРОТИВ ПОЛИТИКИ» · ДОСЛОВНО

«Контекст — нужен; политика — обязательна... Если перехватчиков слишком много, агентный конвейер начинает буксовать. Правило калибровки: политики используем для критичных требований безопасности и соответствия регулированию, а контекст — для соглашений команды»

Эпистемическая честность, местами. Полная версия трижды притормаживает сама себя. О выводе Bain: «это интерпретация результатов исследования, а не жесткий факт». О прогнозах: «направленные сигналы, а не детерминированные утверждения». О цифре 20 человеко-лет: «грубая оценка масштаба сдвига, а не целевая метрика». Три оговорки на 175 страниц. Всё, что дальше, объясняет, где их не хватило.

II

Чем глубже, тем меньше доказательств

В краткой редакции двенадцать источников на 44 полосы. В полной ноль на 175. Обычная логика whiteraper вывернута наизнанку.

Ноль источников на 175 полос

Проверено исчерпывающе: ноль URL, ноль DOI, ноль сносок, ноль концевых ссылок, ни одного раздела «Источники», «Литература» или «References». Документ заканчивается глоссарием.

◆ ЕДИНСТВЕННЫЕ БИБЛИОГРАФИЧЕСКИЕ ЯКОРЯ ВО ВСЁМ ТЕКСТЕ

- **arXiv 2604.14228v1**, ссылка на разбор Claude Code
- Подпись под одной зарплатной таблицей: «Источники: hh.ru / IT Institute Q2 2026, Ведомости / hh.ru... Levels.fyi Q3 2025»

Всё. Больше в документе нет ни одной ссылки, которая позволяла бы связать конкретную цифру с конкретной работой, выборкой и методикой.

Для сравнения, краткая редакция. Двенадцать источников, из них пять Anthropic, один «Внутренние материалы Сбер», один Bain без названия отчёта и года. McKinsey в списке нет вообще, хотя на лендинге airdpc.ru цифра «34–45% прирост продуктивности» подписана «Оценка McKinsey».

Сорок два процента чисел ни на что не опираются

Прошли текст сплошь: около 125 мест с конкретными числами. Из них исключены целевые значения метрик (порог «не ниже 95%» не требует источника, это норматив) и иллюстративные расчёты. Осталось **97 эмпирических утверждений о внешней реальности.**

КАК АТРИБУТИРОВАНО	доля
Назван источник	35%
Безличная формулировка	21%
Атрибуции нет вообще	42%
Внутренние данные Сбера, помечено явно	2%

♦ И ДАЖЕ СРЕДИ «НАЗВАННЫХ»

Конкретная работа с названием и годом указана лишь в **одиннадцати случаях на весь документ**: две работы Bain, две McKinsey, четыре Gartner, PwC, Wobo и один препринт arXiv. В остальных стоит только имя компании: «по данным Anthropic», «McKinsey фиксирует», «Gartner предупреждает».

Границы подсчёта. К «безличной» атрибуции отнесено то, где организация не названа, а источник помянут неопределённо («аналитики», «консультанты», «исследовательские компании»). Одно место может содержать несколько чисел и считается одним локусом. При другом критерии отсекация доли сместятся на единицы процентов, а соотношение «названо / не названо» не сместится.

III

Дорожная карта, по которой нельзя построить план

Сроки заданы в разных главах по-разному. А что от чего зависит, как распределена мощность и где проходит критический путь, не сказано нигде.

Три расхождения в сроках

♦ ПУТЬ L2 → L4

Глава 3.6: L2→L3 за 6–12 мес. + L3→L4 за 12–18 мес.

18–30 месяцев

Глава 7.2: Горизонт 2 (6–18 мес.) целиком покрывает L2→L4

12 месяцев

Расхождение в **1,5–2,5 раза**, и документ его не разрешает. Возможно, переходы предполагались параллельными, но нигде об этом не сказано. Читатель, который составляет план, получает два числа для одного отрезка и никакого способа выбрать.

Окупаемость названа трижды и по-разному: «капитальный, 1,5–3 года» · «окупаются за 12–24 месяца» · «сходится на горизонте 18–24 месяцев». Нижняя граница одной вдвое меньше нижней границы другой.

Платформенная команда делает две работы: «12–18 месяцев для команды из 3–5 инженеров платформы» на постройку IDP, и рядом «платформенная команда из 3–5 человек берёт на себя поддержку 10+ продуктовых команд». Возможно, второе относится к состоянию после запуска. Но документ не говорит, как фазируются работы, как распределена мощность и когда строительство сменяется эксплуатацией.

Пятнадцать компонентов, одно число

Что строится за 12-18 месяцев силами трёх-пяти инженеров, по описанию самого документа:

Kubernetes с custom-операторами · Vault · управляемые PostgreSQL/Redis · S3 с блокировкой объектов, RTO=1 час, RPO=5 минут · граф знаний Neo4j · векторное хранилище Qdrant/Weaviate · ИИ-шлюз LiteLLM · реестр MCP-серверов · координатор A2A · Policy Hook Framework · Specification Registry · Eval Registry · Token Economics Tracker · **Guardian Agent** · Real-time Explainability Layer с отдельной малой LLM и кэшем · PDLC Semantic Layer · Outcome Pricing Engine · Federated Policy Hook Registry

Считать «месяц на компонент» здесь некорректно. Готовые продукты достаточно развернуть и настроить. Что-то надо интегрировать. Что-то писать с нуля. И многое делается параллельно.

♦ ПРЕТЕНЗИЯ В ДРУГОМ

Срок **12-18 месяцев нигде не разложен**. Ни на «поставить готовое / интегрировать / написать своё», ни на людей, ни на зависимости, ни на переход от строительства к эксплуатации. Одно число на пятнадцать разнородных позиций.

При этом про один из компонентов документ пишет отдельно и совсем иначе: для уровней T3-T5 агентов-стражей «зрелых поставщиков на российском рынке нет: **это капитальная инвестиция в собственную разработку**», а про T5 сказано: «только build, высокая трудоёмкость». Как эти две оценки соотносятся, документ не говорит.

Проверка закрывающимся окном

ФИНАЛ ДОКУМЕНТА · ДОСЛОВНО

«Окно входа — 2028-2030 годы — закрывается. **Решение — за СІО и СТО:** зафиксировать текущее состояние и запустить пилот сегодня или войти в новый ландшафт уже догоняющими»

От июля 2026 до начала 2028 проходит восемнадцать месяцев.

Организация, начавшая в день публикации:

- по главе 7 успевает пройти Горизонт 2 и оказаться на L4 ровно к закрытию окна;
- по главе 3.6 к тому же сроку она только выходит из L3, потому что один переход L3→L4 занимает 12-18 месяцев;
- L5 в главе 3.6 срока вообще не имеет: вместо срока стоит условие «только для организаций с платформенной командой и работающим контуром управления».

♦ ВНЕШНЯЯ ПРОВЕРКА · DORA

DORA в материалах о ROI предлагает другой способ считать: строить расчёт на **собственных baseline**, закладывая **J-кривую** (временное падение производительности перед ростом) и получать результат через открытый калькулятор, где каждое допущение видно и меняется.

В сроках Сбера провала нет. Рост идёт линейно с первого горизонта, входных параметров под пересчёт не предусмотрено. Ни J-кривая, ни зафиксированный DORA рост нестабильности доставки в документе не упомянуты, хотя DORA цитируется многократно.

И контрольный вопрос. Сам Сбер (14 000 разработчиков, GigaChat, AIRI, собственная IDE) по краткой редакции находится на уровне 3 из 5. Откуда у организации на сто-пятьсот человек возьмётся путь до L4 за восемнадцать месяцев силами трёх-пяти инженеров, документ не объясняет.

IV

Документ, не сведённый сам с собой

Центральная шкала описана дважды. Шкал зрелости пять вместо заявленных трёх. Одни и те же числа взяты из двух разных источников.

R0-R5 описана дважды, описания не сведены

ОСНОВНАЯ ТАБЛИЦА, ГЛАВА 4.3

ГЛОССАРИЙ, 120 ПОЛОС СПУСТЯ

R0	«Только чтение»	«Агент только предлагает; человек принимает каждое решение»
R2	«Действия в ветке с ревью PR»	«Задачи в рамках сессии; ключевые точки требуют одобрения»
R5	«Эксперимент в песочнице» наименее рискованная среда	«Полностью автономный агент; человек определяет только бизнес-замысел» максимальная ответственность

Формально эти описания можно примирить: «полная автономия в песочнице» и «человек задаёт только бизнес-замысел» способны описывать один сценарий с разных сторон. Нижние ступени примирить труднее: «только чтение» и «агент предлагает» задают разное.

◆ ГЛАВНОЕ

Mapping между тремя осями (**автономия агента, среда исполнения, права записи**) не задан нигде. Читатель, который внедряет шкалу, должен восстановить его сам, а от этого зависит, что именно агенту разрешат делать в проде.

Вторая версия совпадает со шкалой автономии из блога James Read (15.02.2026) примерно на 70–80%, включая тождественную верхнюю ступень «команда агентов» и тот же набор референсных инструментов вплоть до OpenClaw.

Шкал пять, источников два, арифметика третья

ДОКУМЕНТ ПРЕДУПРЕЖДАЕТ САМ · ГЛАВА 3.6

«Обозначение L0–L5 в данном разделе относится к зрелости команды. **Три разные оси...** Смешение осей — типичная ошибка при самооценке зрелости»

И тут же смешивает пять: зрелость команды (гл. 3.6), горизонт задачи и поколение инструментов (гл. 2.5), уровни сжатия контекста (гл. 2.8), зрелость платформы (гlossарий), доля задач × длительность сессии (гл. 6.3). Причём числа между ними расходятся: у L2 «до 30 минут» и «≤15 мин», у L3 «до 2 часов» и «≤1 часа».

♦ ОДНА ПАРА ЧИСЕЛ, ДВА ИСТОЧНИКА

Полоса 3 «отраслевой бенчмарк **DeepSense.ai** на **225 задачах...** с 53,8% до 81,8%»

Полоса 47 «**по данным аналитических агентств...** с 53,8% до 81,8%»

♦ АРИФМЕТИКА, КОТОРАЯ НЕ СХОДИТСЯ С СОСЕДНИМ ПРЕДЛОЖЕНИЕМ

«**Прирост в 5–10 раз** меняет класс достижимых задач: задачи, прежде требовавшие 20 человеко-лет, выполняет команда из 4–6 разработчиков за 3 месяца». Двадцать человеко-лет складываются в 240 человеко-месяцев. Четырнадцать разработчиков за три месяца дают 12–18. Отношение: **13–20 раз**.



Цифры, которые не проверяются

У центрального аргумента документа пять вхождений и ни одного источника. Опорный контраргумент опровергается публичными бенчмарками.

22 п.п. — пять вхождений, ни одного источника

ДОСЛОВНО · ПОЛНАЯ ВЕРСИЯ

«По данным **крупнейших технологических консультантов**, один разработчик с одной моделью в разных средах показывает разброс результата в 22 п. п.»

«Инвестиция в среду работы ИИ-агентов даёт прирост в 22 п.п.; инвестиция в смену модели — не более 3 п.п. **Это измеримая экономика**»

Пять вхождений. Ни в одном не названы ни компания, ни исследование, ни методика, ни размер выборки. При этом утверждение усилено до «это не метафора, а измеримое соотношение».

♦ РЕАЛЬНЫЕ ИЗМЕРЕННЫЕ ВЕЛИЧИНЫ С ОПУБЛИКОВАННОЙ МЕТОДИКОЙ

6 п.п.

Anthropic, «Infrastructure noise»: разрыв между максимально и минимально обеспеченными конфигурациями, $p < 0,01$. На SWE-bench — 1,54 п.п.

×7

Braintrust, 1 781 продакшн-трейс: харнесс объясняет примерно в 7 раз больше дисперсии успеха, чем выбор модели

В публичном поиске такая величина всплывает: на Terminal-Bench 2.0 одна модель под разными харнессами даёт разброс 22,2 п.п. Но это скор бенчмарка обвязки, а не «успешность разработчика», и абсолютные значения не совпадают с приведёнными 45%/23%. Совпадение числа даёт основание для гипотезы, но не доказывает происхождения.

♦ ФОРМУЛИРОВКА, ЗА КОТОРУЮ МОЖНО ОТВЕЧАТЬ

Тезис «среда важнее модели» правдоподобен и частично подтверждается внешними исследованиями. Конкретная цифра 22 п.п. в приведённом виде **не верифицируема**.

«Разрыв между топовыми моделями — 1-3 пункта»

Опровергается публичными данными:

БЕНЧМАРК	РАЗРЫВ
SWE-bench Verified, текущий фронтир (95,0% vs 80,5%)	14,5 п.п.
SWE-bench Verified, полная выборка	22,5 п.п.
SWE-bench Pro, два активных флагмана (69,2% vs 58,6%)	10,6 п.п.
Terminal-Bench 2.0, полный разброс	81,6 п.п.

♦ ОТКУДА, ВЕРОЯТНО, ВЗЯЛИСЬ «1-3»

Anthropic в «Infrastructure noise» рекомендует **не доверять** разнице *меньше 3* п.п. на лидерборде, потому что она может быть шумом инфраструктуры. Порог статистической значимости превратился в утверждение о фактическом разрыве между моделями.

На этом утверждении держится весь стратегический вывод документа: «модель коммодитизируется, инвестируйте в среду». Вывод при этом верен, но не по той причине, которую называет документ.

Две подмены по дороге через документ

♦ 27% МЕНЯЕТ СМЫСЛ

стр. 11 «около 27% **задач** с участием ИИ вообще не были бы сделаны без него» — со ссылкой на Anthropic, корректно

стр. 38 «ROI очевиден: команда производит на 27% больше **ценности** без изменения численности»

стр. 43 «агенты добавляют около 27% **ценности** без роста штата» — уже без атрибуции

Доля задач, которые иначе не делались бы, и прирост производимой ценности меряют разное: задачи, которые раньше не делались, могли не делаться именно потому, что стоили меньше остальных.

♦ 98% — ТРИ ПОДМЕНЫ

Источник реален, и в полной версии процитирован корректно: arXiv 2604.14228, «~1,6%». Но в краткой редакции и на лендинге та же цифра подана как «98% кодовой базы — детерминированная логика» со ссылкой на «анализ ведущих open-source решений».

- строки кода → «инженерная ценность»
- один продукт (декомпилированный бандл Claude Code v2.1.88) → класс решений
- оценка анонимного «community analysis» в непрорецензированном препринте → установленный факт

Прогнозы без имён и метрика, которая не считается

«К 2028 году ручное написание кода сократится вдвое» ·
 «дефекты после развёртывания сократятся на 60%» · «к 2030
 роль product engineer заместит более половины разработчиков» ·
 «без SDD к 2027 отставание на 25% по скорости» · «prompt-to-app
 без governance увеличит число дефектов на 2500%».

Атрибутировано везде одинаково: «по оценкам отраслевых
 аналитиков», «по данным мировых исследовательских
 компаний». Ни одного имени.

♦ CТОR — МЕТРИКА, У КОТОРОЙ НЕТ ЗНАМЕНАТЕЛЯ

Cost-to-outcome ratio объявлена ключевой метрикой экономической
 эффективности. Пороги заданы жёстко: «<1.0 — задача окупилась; <0.3 —
 отличный возврат; >3.0 — повод для ретроспективы».

Формула: общая стоимость токенов / ценность подтверждённого результата.

Знаменатель нигде не операционализирован. Документ признаёт это открытым
 текстом: «команде нужна явная политика оценки стоимости результата».

Пороги абсолютны, а способ вычислить величину оставлен читателю.

Отдельно стоит блок «Ключевые метрики для России»: «100% compliance rate субъектов ПДн»,
 «>99,5% покрытие журнала аудита», «100% операций через отечественные модели». Не указано,
 достигнутый это результат, целевое значение или иллюстрация.

VI

Два недействующих акта в основании

Самая конкретная часть whitepaper. И потому единственная, которую можно проверить не по вкусу, а по документам.

Что обещает RU Compliance Core

ДОСЛОВНО · ГЛАВА 8.3

«Каждая агентная задача, затрагивающая данные, автоматически наследует применимые политики из библиотеки через Policy Hook. **Инженер не должен помнить о 152-ФЗ — Policy Hook напомнит сам**»

ПОЛИТИКА	ОСНОВАНИЕ ПО ДОКУМЕНТУ
PII-Localization Guard	152-ФЗ, ст. 18.1
Data Retention Policy	152-ФЗ, ЦБ РФ 683-П
Processing Purpose Declaration	152-ФЗ, ст. 9
Immutable Audit Trail	ЦБ РФ, ФСТЭК · 5 лет
Domestic Inference Gate	152-ФЗ, Указ №166, ФСТЭК
Sanctioned Entity Check	Росфинмониторинг
Privileged Access Control	ФСТЭК Приказ 17/21 → R3+
Incident Reporting Hook	ФЗ-187, ГосСОПКА

◆ ЧТО В ПОЛНОЙ ВЕРСИИ ИСПРАВЛЕНО — И ЭТО НАДО ПРИЗНАТЬ

Domestic Inference Gate перепривязан с Указа № 250 на **Указ № 166**, и это правильная норма. Пятилетнее хранение аудита откреплено от Указа № 250, где такого требования нет вообще. Сама строка про Указ № 166 воспроизведена с редкой точностью, включая правку Указом № 214 от 07.04.2025.

Проверка по первоисточникам

НОРМА ПО ДОКУМЕНТУ	ЧТО РЕГУЛИРУЕТ НА САМОМ ДЕЛЕ	ВЕРДИКТ
152-ФЗ, ст. 18.1 → локализация	Организационные меры оператора: ответственный, политика, внутренний контроль. Локализация — ч. 5 ст. 18	не та норма
683-П ЦБ → сроки хранения	Противодействие переводам без согласия клиента. Утратило силу 29.03.2025 , заменено 851-П	акт отменён
Приказ ФСТЭК № 17	Защита информации в ГИС. Утратил силу с 01.03.2026 , заменён приказом № 117	акт отменён
Приказы 17/21 → «автономия R3+»	Ни в одном нет понятия уровней автономии — ни агентов, ни систем	своё за норму
«Сертификация Роскомнадзора»	РКН ведёт реестр операторов ПДн. Сертифицируют ФСТЭК и ФСБ	института нет
«Запрет передачи ПДн без согласия»	Ст. 12: уведомительно-разрешительный режим с 01.03.2023, РКН вправе запретить	модель до 2023
Указ № 250	Меры ИБ, персональная ответственность, запрет иностранных СЗИ с 01.01.2025	корректно
Указ № 166	Запрет иностранного ПО на значимых объектах КИИ	корректно

♦ ПОЧЕМУ ЭТО НЕ КРЮЧКОТВОРСТВО

Policy-as-code, построенный на отменённом положении и не той статье закона, **кодифицирует ошибку и делает её невидимой**. Ошибка в правовой ссылке, зашитая в автоматический гейт, масштабируется со скоростью агентного конвейера.

Модель риска неполна — и это работает против них же

Оценка ущерба в документе: «152-ФЗ — до 18 млн рублей за нарушение». **Сама цифра корректна:** 18 млн стоят в ч. 9 ст. 13.11 КоАП верхней границей за повторное нарушение требования о локализации. Эта санкция действует.

Но модель риска ею и исчерпывается. За её пределами остался целый класс составов, введённых 420-ФЗ и действующих с 30 мая 2025 года:

20 млн ₽

за утечку биометрии; до 15 млн — при утечке данных свыше 100 тысяч субъектов

до 500 млн ₽

оборотный штраф 1–3% выручки за повторную утечку. Для банков база — от собственных средств

♦ ТОЧНАЯ ФОРМУЛИРОВКА ПРЕТЕНЗИИ

Это **другое нарушение**, а не переоценка того же. Но именно оно релевантно для агентного конвейера, который обрабатывает персональные данные. Утечка через агента вероятнее, чем повторное нарушение локализации.

И это работает против аргументации самого документа: вся экономика инвестиций в безопасность в главах 6 и 7 построена на оценке ущерба, которая занижает верхнюю границу порядка на два. Реальные цифры сделали бы их же вывод намного сильнее.

Из чего собрана «новая архитектура»

♦ ОГОВОРКА, БЕЗ КОТОРОЙ РАЗДЕЛ ЛЕГКО ОТБИТЬ

Переиспользовать отраслевые термины нормально, это не плагиат. Претензия в другом: документ подан как **новая архитектура**. Между «мы адаптируем подходы Spotify, Gartner, HashiCorp и AWS под российский контур» и «концепция, переосмысляющая жизненный цикл разработки» лежит вся разница.

ТЕРМИН	ПЕРВОИСТОЧНИК	ГОД
Golden Path	Spotify Engineering	2020
Guardian Agents	Gartner, Avivah Litan	06. 2025
Policy as Code	HashiCorp Sentinel	2017
Harness	OpenAI, «Harness engineering»	02. 2026
Context engineering	Lütke → Karpathy → Willison	06. 2025
Spec-Driven Development	GitHub Spec Kit / Amazon Kiro	2025
Уровни автономии	SAE J3016	2014
Шесть мультиагентных паттернов	Anthropic, «Building effective agents» — пять паттернов плюс их же категория, названия и порядок совпадают дословно	12. 2024

Из десяти проверенных терминов атрибутирован ноль. При этом атрибутировать документ умеет: AWS AI-DLC получил отдельный раздел сравнения, Amazon Working Backwards назван, DORA отнесена к Google Cloud, закон Гудхарта и теорема Байеса названы по имени. Просто не делает этого там, где заимствование конституирует его собственную новизну.

Кому и зачем это написано

ПОЛОСА 4 · ДОСЛОВНО

«Документ ориентирован на руководителей высшего звена — **СЕО, СТО, СЮ, СISO, САЮ** и других лидеров цифровой трансформации»

Инженер здесь не аудитория. Он объект, о котором принимают решение.

Продают не продукт. GigaCode и GigaIDE не упомянуты ни разу ни в одной из версий, все продуктовые кейсы майской редакции вычищены. Вместо этого идёт атака на саму модель лицензирования: «налог на поставщика», «\$1–3 млн в год для компании из 100 человек», «за эти же деньги можно построить собственную платформу».

♦ **МОТИВ: ЧЕТЫРЕ ВЕРСИИ, НИ ОДНА НЕ ДОКАЗАНА**

- **Заявка на отраслевую рамку.** На неё работают три языка лендинга; прямых свидетельств нет
- **Создание рынка под свою позицию.** За версию говорит атака на лицензии, но коммерческого предложения в документе нет
- **Внутреннее лоббирование бюджета.** Самое скромное объяснение и, возможно, самое точное
- **Персональная позиция автора.** Подпись стоит на каждой из 175 полос, плюс два федеральных выступления

Разбор на мотиве не строится и в нём не нуждается: всё в частях II–VI проверяемо независимо от того, зачем документ написан. Отдельная деталь жанра: публикация в «Коммерсанте» вышла с параметром `erid=`, а это обязательная маркировка материала как интернет-рекламы.

Пять вопросов для переговорной

Когда прозвучит «давайте делать как в Сбере»:

1 **Воспроизведите цифру 22 п.п.**

Какое исследование, какая выборка, что считалось «успешностью»? Если ответ — «данные крупнейших консультантов», решение о капзатратах опирается на утверждение, которое нельзя ни подтвердить, ни опровергнуть

2 **Соберите из этого план.**

Глава 3.6 даёт 18–30 месяцев, глава 7.2 — двенадцать. Какая цифра идёт в план? Что параллелится? Где критический путь?

3 **Посчитайте платформенную команду.**

Пятнадцать компонентов, включая Guardian Agent, который сам документ называет отдельной капитальной инвестицией. Три-пять инженеров. И те же люди параллельно поддерживают десять команд

4 **Проверьте правовые ссылки до кодирования в хуки.**

Почему локализация обоснована статьёй про организационные меры? Почему срок хранения выводится из положения, отменённого в марте 2025?

5 **Покажите свой baseline и свои допущения.**

У DORA под ROI лежит открытый калькулятор, где каждое допущение видно и меняется. Под цифрами whitepaper не лежит ничего, что можно пересчитать под вашу компанию

Ни один из этих вопросов не требует спорить с тезисом «среда важнее модели». Он верен. Спор о том, можно ли по этому документу планировать бюджет.

Брать · проверять · не принимать на веру

БРАТЬ	ПРОВЕРЯТЬ ПЕРЕД ИСПОЛЬЗОВАНИЕМ	НЕ ПРИНИМАТЬ НА ВЕРУ
Формула эффективный_R с восемью факторами эскалации	Матрицу Guardian Agents T1-T5 — применимость продуктов не проверялась	Цифру 22 п.п. — источник не назван ни в одном из пяти вхождений
Структуру Evidence Bundle: восемь обязательных элементов	Метрики главы 6 — формулы есть, но CTOR не считается без своей политики оценки	Сроки L2→L4 — 12 месяцев против 18-30 в разных главах
Cost Circuit Breakers: задача, сессия, горизонт	Уровни R0-R5 — свести два описания и задать mapping самостоятельно	Экономику IDP — без обоснования и без декомпозиции
Различение «контекст против политики»	Каждую ссылку RU Compliance Core — два акта из пяти недействительны	Прогнозы «к 2028 / 2030» — сам документ называет их направленными сигналами
Финальные три шага: baseline по DORA, пилот на 2-3 командах, AI Champions	Матрицу зрелости L0-L5 — в документе пять разных шкал под этим именем	«-60% дефектов» — прогноз без методики в якорных цифрах

Последняя строка левой колонки повторяет ровно то, что рекомендует сам документ в финале. Единственная его рекомендация, которая не требует веры в его же цифры.

Что не проверено

- **Признаки машинной генерации.** Разбор фиксирует внутреннюю несогласованность документа. Как именно текст создавался, отсюда не следует и здесь не утверждается.
- **Сплошной визуальный просмотр всех 175 полос не проводился.** Растровых изображений нет, это установлено технически; векторных диаграмм не нашлось, но искали выборочно.
- **Метаданные PDF** (автор файла не совпадает с автором по подписи) приведены как факт; для корпоративных документов это обычная практика.
- **Продукты матрицы T1-T5** (GitVerse plugin API, SonarQube CE, PT Application Inspector, Solar appScreener) на заявленную применимость не проверялись.
- **Источник цифры 22 п.п. не установлен.** Совпадение с Terminal-Bench 2.0 остаётся гипотезой.
- **Экономика IDP не опровергнута.** Она просто ничем не подтверждена и не сходится сама с собой.
- **Правовые вопросы, которые не решаются по открытым источникам:** считается ли «аномальное поведение агента» компьютерным инцидентом по 187-ФЗ; надо ли уведомлять РКН при инференсе через зарубежный endpoint без хранения.
- **Мотив** реконструирован по структуре документа, это не установленный факт.

Двадцать одна правка самого разбора

Разбор прошёл семь кругов независимого ревью. Тринадцать фактических ошибок — все свои, все перепроверены по источнику. Журнал оставлен намеренно: текст, обвиняющий в невоспроизводимости, обязан показывать собственную историю проверок.

v1→v2 «Сбер упомянут только в колонтитуле» — неверно, 42 колонтитула плюс 4 упоминания в теле. Страница оговорки указана неверно

v2→v3 **Документу приписана оговорка, которой в том месте нет.** Реальная цитата, переставленная в другой контекст — ровно тот тип ошибки, за который разбор критикует whitepaper

v3→v4 Мишень раздела о заимствованиях перенесена с факта заимствования на подачу компиляции как новой архитектуры

v4→v5 **DORA описана неверно:** «ROI 39%» подан как результат исследования, тогда как это выход калькулятора на конкретных допущениях

v5→v6 Прочитана полная версия. **Снята** претензия о заимствовании AWS AI-DLC — там есть честный раздел сравнения. Снята часть правовых претензий

v6→v7 **Ошибка с DORA вернулась при переписывании.** Плюс снято «занижение штрафов в 25 раз» — сравнивались разные составы

v7→v8 Устранены внутренние несоответствия самого разбора; из вердикта убрана догадка о людях

v8→v9 Раскрыта методика подсчёта 97 утверждений; смягчено утверждение об оригинальности формулы

Всё проверяемо

ДОКУМЕНТЫ

Whitepaper, краткая —

aipdlc.ru/documents/ru/whitepaper_short_ru.pdf

Whitepaper, полная —

aipdlc.ru/documents/ru/whitepaper_full_ru.pdf

Лендинг — aipdlc.ru/ru

Коммерсантъ — kommersant.ru/doc/8798345 (URL содержит erid=)

ПРАВО

152-ФЗ, ст. 18 ч. 5 — локализация ПДн

152-ФЗ, ст. 18.1 — меры оператора

152-ФЗ, ст. 12 — трансграничная передача

Положение ЦБ 851-П от 30.01.2025 — cbr.ru

Указ № 250 от 01.05.2022 — pravo.gov.ru

Указ № 166 от 30.03.2022

Приказы ФСТЭК 17, 21, 117

Ст. 13.11 КоАП и 420-ФЗ от 30.11.2024

ФАКТ-ЧЕК ЦИФР

Anthropic — «Infrastructure noise», «Effective context engineering», multi-agent research system

arXiv 2604.14228 — «Dive into Claude Code»

SWE-bench Verified — llm-stats.com

SWE-bench Pro — labs.scale.com

Terminal-Bench 2.0 — tbench.ai

METR — RCT на опытных разработчиках, 07.2025

РОДОСЛОВНАЯ

Spotify — Golden Paths, 08.2020

Gartner — Guardian Agents, 06.2025

HashiCorp — Sentinel, Policy as Code, 09.2017

OpenAI — Harness engineering, 02.2026

Anthropic — Building effective agents, 12.2024

AWS — AI-DLC, Kiro

GitHub — Spec Kit, 09.2025

SAE J3016 · Team Topologies · James Read, шкала L0-L5, 02.2026

КОНТЕКСТ

DORA 2025 — State of AI-assisted Software Development

DORA 2026 — ROI report и открытый калькулятор

Thoughtworks Radar Vol. 34 — SDD в кольце Assess

Хабр — разбор майской редакции, С. Попов, 23.05.2026

Спасибо, что дочитал.

Этот разбор — часть канала «Под капотом»: разборы ИИ-индустрии без хайпа, с источниками, датировками и пометкой «не проверено» там, где данных нет. Если зашло, забирай продолжение.

TELEGRAM · КАНАЛ

Под капотом

t.me/gorilla_under_hood ↗



Лого
канала

Автор

Илья Утов · AI Frontier

LinkedIn

linkedin.com/in/ilyautov

Что делаем

Строим и отгружаем ИИ-системы. Эффект, а не слайды.