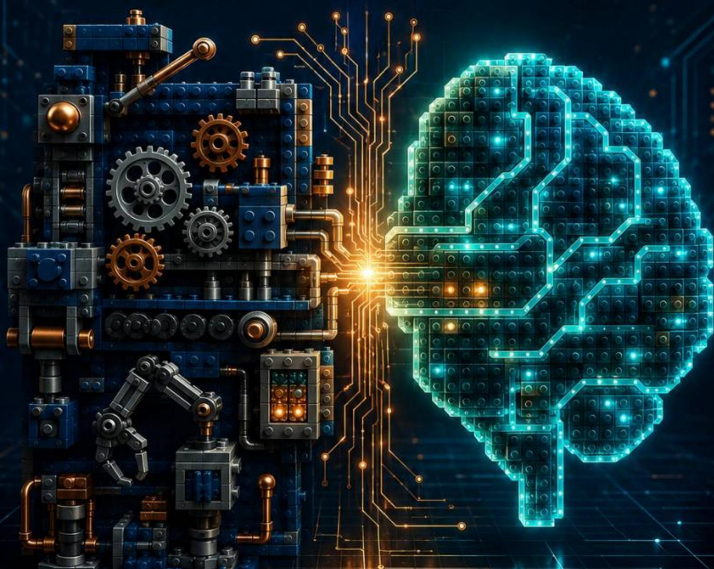


ПОЛЕВОЙ ГИД ПО АРХИТЕКТУРЕ СОВРЕМЕННОГО AI-АГЕНТА

Десять исследований [arXiv](#),
которые изменят ваш [AI-продукт](#) в 2026 году



Павел Новопольцев

Полевой гид по архитектуре

современного AI-агента

<https://litres.ru/74300188>

SelfPub; 2026

Аннотация

Почему один и тот же искусственный интеллект в руках одной команды работает, а в руках другой нет? Не потому что одна команда покупает более дорогую модель и не потому что у нее есть секретные слова для промптов. Дело в обязанности.

Обязанка - это всё, что окружает языковую модель в работающей системе: системный промпт, набор инструментов, способ хранения контекста между сессиями, цикл обратной связи, в котором модель проверяет собственные ответы и пересобирает стратегию

Десять свежих исследований об архитектуре AI-агентов, переведённые в практику. JSON умирает, Python-стабы побеждают.

Промпт с «подумай глубже» сжигает бюджет в 18 раз.

Полная история диалога ломает агента.

Память, которая помнит нужное, повышает успех с 79 до 96 процентов.

Книга для тех, кто внедряет AI-агентов в продакшн и хочет перестать палить деньги.

Содержание

Введение	5
Что вы получите	8
Как читать эту книгу	9
Чего в этой книге нет	10
Глава 1. Горький урок 2026 модель это тридцать процентов, обвязка семьдесят	11
Почему это важно	14
Цифры, которые изменяют ваш подход	16
Что считать обвязкой	17
Что вы узнаете в следующих главах	19
Что сделать прямо сейчас	20
Глава 2. Конец JSON как Python-стабы обходят классический function calling	21
Два способа позвать инструмент	23
Цифры, которые изменяют ваш подход к инструментам	25
Конец ознакомительного фрагмента.	26

Павел Новопольцев

Полевой гид

по архитектуре

современного AI-агента

Введение

Эта книга о том, почему один и тот же искусственный интеллект в руках одной команды работает, а в руках другой нет. Не потому что одна команда покупает более дорогую модель. Не потому что у нее лучше промпт-инженеры. И уж точно не потому, что она знает секретные слова, которые заставляют нейросеть «стараться сильнее».

Дело в обвязке.

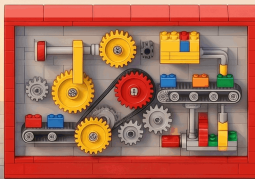
Обвязка - это все, что окружает языковую модель в работающей системе: системный промпт, набор инструментов, способ хранения контекста между сессиями, цикл обратной связи, в котором модель проверяет собственные ответы и пересобирает стратегию. Модель - это мозг. Обвязка это руки, инструменты, память и рабочий стол. И, как показывают свежие исследования, разница между «способной моделью с плохой обвязкой» и «средней моделью с хорошей» в 2026

году огромна.

Эту книгу я писал для тех, кто внедряет AI-агентов в реальные продукты. Для инженеров, которые видят, что агент «вроде работает», но жжёт бюджет и проваливает простые задачи. Для продакт-менеджеров, которые не понимают, почему замена модели на более дорогую не дала прироста. Для основателей небольших студий, которые хотят разговаривать с разработчиками на одном языке, когда речь идет об архитектуре агентской системы.

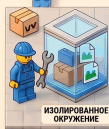
Здесь нет кода, который нельзя запустить, и нет формул, которые нельзя объяснить. Здесь есть десять исследований, отобранных по одному критерию: применимо ли это завтра в вашем продукте. Каждое из них опубликовано на arXiv в июле или августе 2026 года, и каждое меняет что-то в том, как нужно строить агентов.

AgentExecutor и E3: Проектирование Сверхэффективных ИИ-Агентов



ТРЕХФАЗНЫЙ ПРОЦЕСС AgentExecutor

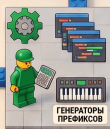
СТРАТЕГИЯ E3: ЭКОНОМНОЕ МЫШЛЕНИЕ



ФАЗА 1: EnvAgent — Подготовка среды
Агент минимизирует изолированное окружение, устраняет избыточности (проц и т.) и отсуствует неадекватные ресурсы.



ФАЗА 2: RofAgent — Итеративное уточнение
Данная фаза исследовании путей выполнения и использование «дружелюбной» обратной связи в состоянии контекста (prompting).



ФАЗА 3: EvoAgent — Синтез генераторов
Синтезирует программно-генератор, который систематически паразитирует прочие варианты профиссов для достижения максимального покрытия.



ESTIMATE (Оценка) — Определение сложности
Агент определяет необходимый объем информации перед началом работы, чтобы стратегия «максимального контекста» для простых задач.

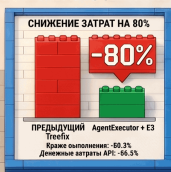


EXECUTE (Выполнение) — Минимальный путь
Запрос самого короткого неадекватного пути выполнения. Это позволяет сэкономить до 91% топлива на простых задачах.



EXPAND (Расширение) — Адаптивный роот
Область поиска и контекст расширяются только тогда, когда проверено минимальное количество собранной информации.

РЕЗУЛЬТАТЫ И ЭФФЕКТИВНОСТЬ



ACRR — Коэффициент когнитивной избыточности
Матрица, измеряющая объем «бесполезной» работы агента. Методы E3 и AgentExecutor стремятся свести этот показатель к минимуму.

Что вы получите

Книга состоит из трех частей.

В первой части мы разберемся, почему обвязка важнее модели, и увидим два конкретных способа, которыми плохая обвязка буквально сжигает деньги. Это главы о JSON-вызовах инструментов, которые проигрывают Python-стабам, и о промптах, которые заставляют агента работать в семь раз больше, чем нужно.

Во второй части перейдем к архитектуре. Научим агента понимать, насколько задача сложная, и не читать весь репозиторий ради однострочной правки. Построим память, которая помнит нужное и не забывает лишнее. Защитимся от ситуации, когда старая история диалога ломает агента, и научимся делать retrieval так, чтобы контекст не раздувался.

В третьей части поговорим о циклах, которые улучшают сами себя. Агент, который переписывает собственную обвязку в рантайме. Система из нескольких агентов, которые эволюционируют общую обвязку на своих данных и делятся находками. И, наконец, метрики, по которым вообще можно понять, что обвязка стала лучше, а не просто «вроде работает».

Как читать эту книгу

Книгу можно читать последовательно, от введения до заключения. Главы выстроены так, что каждая опирается на предыдущую и готовит следующую. Но если вы опытный инженер и вам нужна конкретная тема, можно начать с нее каждая глава самодостаточна, и в конце есть мостик к следующей, чтобы вы не потерялись.

Если вы читаете книгу, чтобы внедрить что-то конкретное в понедельник утром, в каждой главе есть блок «Что сделать прямо сейчас» с конкретными шагами. Если вы читаете, чтобы понять, как устроено современное AI-внедрение, эти блоки можно пропустить.

В конце книги есть заключение, которое называется «Что делать в понедельник утром». В нем собран один список из десяти шагов, ранжированных по приоритету. Если у вас есть пятнадцать минут и работающий AI-агент, начните отсюда.

Чего в этой книге нет

Здесь нет гимна возможностям LLM. Здесь нет новостей про очередную модель, обогнавшую предыдущую. Здесь нет обещаний, что через год агенты заменят всех разработчиков.

Здесь есть трезвая инженерная работа. Архитектура. Метрики. Циклы обратной связи. Экономика токенов. И десять конкретных исследований, каждое из которых улучшит вашу систему, если вы внедрите его результаты.

Давайте начнем.

Глава 1. Горький урок 2026 модель это тридцать процентов, обязанка семьдесят

В августе 2026 года исследователи из Scale AI опубликовали работу под названием HarnessOpt-Bench.

Они поставили простой, но неприятный вопрос если дать разным языковым моделям одинаковую задачу и одинаковую обязанку, кто выиграет?

А если дать одинаковую модель, но разные обязанности, кто выиграет тогда?

Весы сложности: Модель vs. Обяззка

Модель vs. Обяззка: Скрытая архитектура ИИ-агентов

Проблема когнитивной избыточности



COMPARISON: Модель («Мозг»)
Ядро логики (LLM).
Маленький «кватернионный мозг» в стальной сфере – мощный, но изолированный компонент, отвечающий за предоставление следующего токена.



COMPARISON: Обяззка («Harness»)
Двигатель исполнения.
Сложная машина с шестернями, трубами и инструментами, которая управляет посылкой файлов, вызовом API и проверкой результатов. На исках она значительно тяжелее модели.



KEY FINDING: Влияние на успех
Обяззка определяет результат.
Оптимизация обяззка (промпты, контроль потока, инструменты) дает больший прирост производительности, чем замена самой модели.



$$ACRR = \frac{C_{\text{факт}} - C_{\text{мин}}}{C_{\text{мин}}}$$

Формула петера: Показатель того, во сколько раз фактические затраты агента (токены, арены, файлы) превышают минимально необходимые для решения задачи.

SUPPORTING FACT: Парадокс простоты
Избыточность выше на простых задачах. Агенты часто используют стратегии максимального контекста, перечисляя знакомые файлы для триксимальных правок, что увеличивает затраты на 400% и более.



Решение: Фреймворк E3 (Estimate, Execute, Expand)



PROCESS STEP: 1. Оценка (Estimate)
Определение начальной рабочей точки. Агент заранее оценивает сложность задачи и необходимый объем информации перед тратой бюджета.

PROCESS STEP: 2. Выполнение (Execute)
Путь минимального сопротивления. Запуск решения с использованием только тех инструментов и файлов, которые были определены на этапе оценки как критические.



PROCESS STEP: 2. Выполнение (Execute)
Итеративный рост области поиска. Контекст расширяется только в случае верификации и файлов, которые имеют до 85% затрат при критических.



PROCESS STEP: 3. Расширение (Expand)
Итеративный рост области поиска. Контекст расширяется только в случае провала верификации, что позволяет экономить до 85% затрат при 100% успехе.

AgentExecutor и автоматическая эволюция

KEY FINDING: Эффективность AgentExecutor
Покрываемость кода до 94% при снижении затрат на 55%.
Трехфазный дизайн: подготовка среды, динамическое исследование с сокращением контекста и синтеза программ для генерации префиксов.



EXAMPLE: EvolveNet: Коллективное самосовершенствование
Обмен опытом без передачи данных. Агенты в разных средах локально улучшают свои «обяззки», а затем обобщают удачные программные адаптации в единый глобальный стандарт.

Data Table: Сравнение эффективности политик (на базе MSE-Bench)

Политика	Успех (Solved)	Стоимость (C)	Избыточность (ACRR)
Max Context First	100%	122.9	12.90
Adaptive Retrieval	100%	22.1	1.21
E3 (Наш подход)	100%	18.6	0.55
Fixed ReAct	66.9%	17.2	1.29

Ответ оказался таким: обязанка важнее.

В эксперименте участвовали пять frontier-моделей: Claude Opus 5, Claude Sonnet 5, Kimi K3, GPT-5.6-Sol и GPT-5.6-Terra. Все они решили четыре разные задачи GAIA (многошаговое исследование с веб-инструментами), OfficeQA Pro (корпоративное рассуждение над документами).

ми), BrowseComp-Plus (*глубокий web-research*) и Terminal-Bench 2.0 (задачи command-line).

Каждая модель работала в двух режимах: под общей coding-обвязкой и под своей нативной обвязкой. Сто одиннадцатая зачетных прогонов. И результат, который стоит запомнить.

Когда меняли модель, эффект на качество был огромным.

Claude Opus 5 стабильно превосходил остальных. Когда меняли обвязку, эффект тоже был, но меньше. Но и это критическая оговорка оптимизация обвязки, выполненная правильно, давала прирост, **сопоставимый** с разницей между лучшей и худшей моделью.

Переведем на человеческий. Если у вас есть средняя модель и хорошая обвязка, вы обгоняете тех, у кого топовая модель и плохая обвязка. Если у вас топовая модель и плохая обвязка, вы проигрываете.

Почему это важно

Последние пять лет мы жили в парадигме «купить более умную модель». Бенчмарки росли, новости выходили, маркетологи продавали. Но между «модель хорошо решает задачу в лаборатории» и «агент хорошо решает задачу в продакшне» — пропасть. Эта пропасть и заполняется обвязкой.

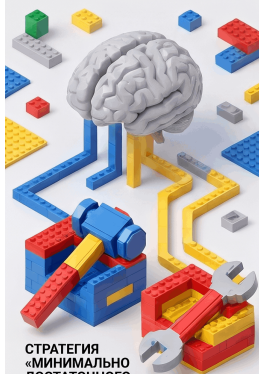
Обвязка делает три вещи, которые модель сама по себе делать не умеет.

Во-первых, обвязка превращает абстрактные способности модели в конкретные действия. Модель умеет рассуждать. Но чтобы рассуждение превратилось в ответ клиенту, нужны: правильный формат вывода, проверка через инструмент, переспрос при неоднозначности, логирование для отладки. Все это — обвязка.

Во-вторых, обвязка ограничивает то, что модель может сломать. Модель умеет генерировать SQL. Но чтобы SQL не угробил базу клиента, нужны белый список таблиц, проверка синтаксиса, тестовый прогон на копии, откат при ошибке. Все это обвязка.

В-третьих, обвязка создает экономику. Модель стоит одинаково на любом запросе. Но архитектура вызова может сделать этот запрос в десять раз дешевле или в десять раз дороже. И выбрать дешевый способ — это инженерное решение, а не магия модели.

ЭФФЕКТИВНОСТЬ ИИ-АГЕНТОВ: КАК ИЗБЕЖАТЬ «ИЗБЫТОЧНОГО МЫШЛЕНИЯ»



СТРАТЕГИЯ «МИНИМАЛЬНО ДОСТАТОЧНОГО ИСПОЛНЕНИЯ»

Агент сначала оценивает сложность задачи и строит кратчайший путь, прежде чем тратить бюджет.



ИНСТРУМЕНТАЛЬНЫЙ СТЕК АГЕНТА

Использование Python-скриптов и Bash-команд для активного взаимодействия с кодом.



ПРОГРЕССИВНОЕ РАСШИРЕНИЕ (EXPAND)

Область поиска и контекст расширяются только в случае неудачи начального плана.



ВЫБОРОЧНАЯ ПАМЯТЬ (SELECTIVE MEMORY)

Система сохраняет спецификации задач и схем данных, но отображает лишние логи рассуждений.



ЗАЩИТА ОТ «ЗАСОРЕНИЯ» ИСТОРИИ

Удаление нерелевантных шагов предотвращает деградацию качества ответов при длительных сессиях.



ACRR: КОЭФИЦИЕНТ ИЗБЫТОЧНОСТИ

Избыточность мышления наиболее высока на самых простых задачах.



СНИЖЕНИЕ ЗАТРАТ НА 85%

Фреймворк E3 сокращает расходы на токены и время выполнения без потери успеха.

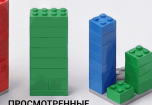


MAX-CONTEXT- FIRST



УСПЕШНОСТЬ: 100%

УСПЕШНОСТЬ: 100%



MAX-CONTEXT-FIRST

ПРОСМОТРЕННЫЕ
ФАЙЛЫ:
8%

СТРАТЕГИЯ E3
(OURS)

ЭФФЕКТИВНОСТЬ AGENTEXECUTOR

Покрытие кода увеличивается до 94% при сокращении числа вызовов LLM вдвое.



Цифры, которые изменяют ваш подход

Возьмите эти числа и запомните. Они из разных исследований, но все сходятся в одном.

Claude Opus 5 под плохой обвязкой проигрывает Claude Sonnet 5 под хорошей обвязке. Это не исключение. Это правило.

Хорошая обвязка может сократить расход токенов в десять раз без потери качества. Или улучшить качество в два раза при том же расходе. Или и то и другое одновременно.

Пять-тридцать раз -это разброс стоимости одной успешной задачи в зависимости от обвязки, которую вы выбрали. Те же модели, те же задачи — разная обвязка, разная экономика.

Это значит, что бюджет на AI-продукт нужно планировать не как «закупка API», а как «инвестиции в обвязку». Если вы тратите сто тысяч рублей в месяц на API и не тратите ничего на обвязку, вы делаете ровно наоборот.

Что считать обязанностью

Чтобы дальше в книге не было путаницы, давайте зафиксируем, что входит в это понятие.

Обязанность - это программа, которая вызывает языковую модель, передает ей контекст, получает ответ, проверяет его и либо возвращает пользователю, либо переспрашивает модель. Звучит просто, но в реальности обязанность — это десятки и сотни решений.

- Какой системный промпт использовать.
- Какие инструменты дать модели.
- В каком формате она должна вернуть ответ.
- Как обрабатывать ошибки.
- Как хранить историю диалога.
- Как выбирать между несколькими моделями.
- Как считать бюджет. Как логировать.
- Как тестировать.
- Как версионировать.
- Как катить в продакшн.
- Как откатывать, если сломалось.

Каждое из этих решений - это строка в счете за API.

Каждое решение -это потенциальный баг. Каждое решение - это точка, где ваш агент либо работает, либо нет.

Большинство команд, которые внедряют AI-агентов, тратят девяносто процентов времени на выбор модели и десять

процентов на обязательку. Это перевернутая пропорция. Правильная - наоборот.

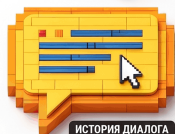
ПРОБЛЕМА «СВЕРХОБДУМАННОСТИ»: ПОЧЕМУ ИИ-АГЕНТЫ ТРАТЯТ ВАШИ РЕСУРСЫ?

Риски и коренные причины неэффективности ИИ-агентов в инженерии и коде

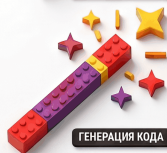


КОГНИТИВНАЯ
ИЗБЫТОЧНОСТЬ

Коэффициент AC RR:
избыточности (ACRR)
5X БОЛЬШЕ РЕСУРСОВ



ИСТОРИЯ ДИАЛОГА



ГЕНЕРАЦИЯ КОДА

ИЗБЫТОЧНЫЙ КОНТЕКСТ
СНИЖАЕТ КАЧЕСТВО

Финансовые и
временные потери

Замедление
выполнения на 80%



ОГРАНИЧЕННОЕ
ПРОСТРАНСТВО ДЕЙСТВИЙ

ОТСУТСТВИЕ
МЕТА-СУЖДЕНИЯ

РАЗМЫТИЕ
ВНИМАНИЯ

СТРАТЕГИЯ
«МАКСИМАЛЬНЫЙ
КОНТЕКСТ»

Что вы узнаете в следующих главах

В следующей главе мы разберем одну конкретную часть обвязки вызовов инструментов.

И увидим, что JSON-вызов, который сейчас использует большинство, проигрывает Python-стабам в одиннадцать из четырнадцати моделей.

Это не «иногда лучше» это «почти всегда лучше».

В третьей главе мы посмотрим, как плохо написанный промпт может заставить агента тратить в восемнадцать раз больше токенов без улучшения качества.

И как одно предложение в системном промпте это лечит.

Вторая часть книги - про архитектуру думающего агента. Мы научим агента понимать, насколько задача сложная, и не делать лишней работы. Построим память, которая помнит нужное. Защитимся от ситуации, когда старая история диалога ломает агента.

Третья часть - про циклы, которые улучшают сами себя. Агент, который переписывает собственную обвязку. Система из нескольких агентов, которые эволюционируют общую обвязку. Метрики, по которым это вообще можно измерить.

Все это не теория, все это результаты, опубликованные в 2026 году, с цифрами, бенчмарками и кодом, который можно скачать и запустить.

Все это применимо к вашему агенту в понедельник утром.

Что сделать прямо сейчас

Прежде чем мы двинемся дальше, одна конкретная вещь, которую вы можете сделать сегодня.

Откройте свой текущий проект с AI-агентом.

Найдите место, где вы вызываете языковую модель. Посмотрите, сколько строк кода вокруг этого вызова. Это и есть ваша текущая обвязка. Если там меньше пятидесяти строк, у вас, скорее всего, нет обвязки. У вас есть тонкий слой поверх API.

Это нормально, если вы только начали. Это проблема, если вы внедряете в продакшн. Потому что продакшн - это про обвязку. Модель дает вам базовую способность. Обвязка превращает способность в продукт.

Запомните этот тезис. Мы будем возвращаться к нему в каждой главе.

Дальше - глава 2, в которой мы разберемся, почему классический JSON-вызов инструментов уже проиграл, и что с этим делать.

Глава 2. Конец JSON как Python-стабы обходят классический function calling

В шестом августа 2026 года вышла работа под названием The Bitter Lesson of Tool Calling. Название говорящее.

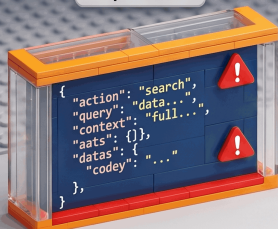
Авторы Ишан Патель, Сахил Сен, Элиас Лумер и Вамсе Кумар Суббия из PricewaterhouseCoopers сравнили два способа дать модели доступ к инструментам.

Результат оказался настолько убедительным, что теперь, спустя несколько недель после публикации, любой инженер, который продолжает использовать старый подход, просто теряет деньги клиента.

Эволюция ИИ-агентов: От жесткого JSON к программному коду

Сравнение подходов: Старая школа vs. Новая эра

Старая школа



Многоходовые задержки

Каждый шаг требует отдельного цикла работы LLM, что замедляет сложные цепочки действий.



Чрезмерный расход токенов

13-кратная когнитивная избыточность
Традиционные агенты переходят весь контекст, увеличивая расходы на 90% без пользы для результата.

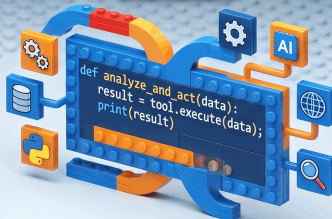


Обрыв вызовов при N > 70

Нативные JSON-парадигмы часто теряют данные при высокой параллельности задач.



Новая эра



Код как универсальный интерфейс

Использование Python-заглушек повышает точность длинных цепочек на 18.8%.



Стратегия E3 (Оценка → Выполнение → Расширение)

Экономия затрат на 85%

Агент сначала оценивает минимально необходимый контекст и расширяет его только при сбоях.



Фреймворк AgentExecutor

Автономная подготовка окружения и эволюция префиксов сокращают время работы на 80%.



Эффективность новых систем

Традиционный подход
(Treefix/JSON)

Время выполнения

Новый подход
(AgentExecutor/E3)

100% (базовое)

-80.3%

Стоимость API

100% (базовое)

-56.6%

Охват кода

79-86%

90-94%

Два способа позвать инструмент

Первый способ, который сейчас используют все, называется JSON tool calling. Работает так: вы передаете модели JSON-схему, описывающую каждый инструмент. Модель возвращает структурированный JSON-объект, в котором указано имя функции и аргументы.

Ваш код ловит этот объект, вызывает функцию, возвращает результат обратно в модель. На каждом вызове отдельный ход инференса.

Если нужно вызвать пять инструментов подряд, это пять ходов. Если нужно распараллелить семь вызовов, модель должна эмитировать семь JSON-объектов за один ответ, и на больших числах она начинает их терять.

Второй способ - programmatic tool calling, или PTC.

Здесь вы передаете модели не схему, а исходный код Python-модуля со стабами функций.

Стаб - это функция с типизированной сигнатурой, которая при вызове делает ровно то, что вы хотите. Модель пишет обычный Python-скрипт, импортирует нужные стабы, вызывает их с аргументами, обрабатывает результаты.

На все семь параллельных вызовов один скрипт, один subprocess, один ход инференса.

Звучит как косметическое отличие. На практике это разница в одиннадцать из четырнадцати протестированных мо-

делей.

Цифры, которые изменяют ваш подход к инструментам

Авторы протестировали четырнадцать языковых моделей на бенчмарке BFCL v4, который специально создан для оценки **function calling**.

Среди моделей - пять моделей семейства Claude (включая Opus 5 и Sonnet 5), несколько поколений GPT, модели от Google и DeepSeek.

Каждая модель запускалась в двух режимах нативный **JSON tool calling** и PTC с Python-стабами.

Одиннадцать из четырнадцати моделей в режиме PTC показали результаты не хуже, чем в режиме JSON. Семейство GPT-5.6 дало прибавку в десять и шесть десятых процента. Тринадцать из четырнадцати моделей не уступали PTC в режиме параллельного fan-out то есть когда нужно одновременно вызвать много инструментов.

Конец ознакомительного фрагмента.

Текст предоставлен ООО «Литрес».

Прочитайте эту книгу целиком, [купив полную легальную версию](#) на Литрес.

Безопасно оплатить книгу можно банковской картой Visa, MasterCard, Maestro, со счета мобильного телефона, с платежного терминала, в салоне МТС или Связной, через PayPal, WebMoney, Яндекс.Деньги, QIWI Кошелек, бонусными картами или другим удобным Вам способом.