

Метод доверенной контекстуализации Model Context Protocol для интеллектуальных агентов маркетинговой аналитики

Г.А. Нижниченко¹, В.А. Павлов²

¹Московский финансово-юридический университет МФЮА, Москва, Россия

²Национальный исследовательский Московский государственный строительный университет, Москва, Россия

Аннотация – Предложен метод доверенной контекстуализации MCP-Trust для интеллектуальных агентов маркетинговой аналитики, решающий задачу воспроизводимого подключения больших языковых моделей к неоднородным корпоративным и открытым источникам данных. В отличие от обычного применения Model Context Protocol как транспортного слоя, метод вводит формализованный контекстный пакет; ранжирование MCP-серверов по релевантности, свежести, надёжности, доступности и стоимости; доказательный граф источников; механизм согласования противоречивых данных и журнал трассировки для повторного воспроизведения результата. Разработан математический аппарат выбора серверов и оценки доверия к аналитическому выводу. Проведён имитационный эксперимент на 180 синтетических аналитических запросах с 12 виртуальными MCP-серверами. По сравнению с вариантом без MCP средняя абсолютная процентная ошибка MAPE снижена с 15,93 до 8,28 %, покрытие результата источниками увеличено с 0,233 до 0,971, коэффициент воспроизводимости – с 0,46 до 0,99. Интегральная метрика эффективности MCP-Trust составила 0,270 против 0,228 у SAMIA и 0,190 у статической MCP-схемы. Полученные результаты подтверждают перспективность использования доверенной контекстуализации для построения интеллектуальных систем исследования рынка.

Ключевые слова – Model Context Protocol, интеллектуальные агенты, маркетинговая аналитика, доказательная верификация, провенанс данных.

ВВЕДЕНИЕ

Интеллектуальные агенты на основе больших языковых моделей всё чаще рассматриваются как средство автоматизации прикладной аналитики: они способны интерпретировать запрос пользователя, вызывать внешние инструменты, агрегировать сведения из нескольких источников и формировать текстовые либо табличные отчёты. Для маркетинговой аналитики это особенно важно, поскольку исследование рынка опирается на разнородные данные: отчёты продаж, CRM-события, сегменты клиентов, публичные показатели конкурентов, данные социальных медиа,

экспертные описания отраслей и открытые статистические наборы.

Вместе с тем прямое подключение языковой модели к источникам данных не устраняет ключевые риски: фрагментарность контекста, отсутствие воспроизводимости ответа, невозможность проверить происхождение отдельных чисел, конфликт между устаревшими и актуальными источниками, а также рост стоимости обращения к внешним инструментам. Появление Model Context Protocol (MCP) создаёт основу для стандартизированного взаимодействия приложений больших языковых моделей с внешними ресурсами и инструментами [1, 2]. Однако сам факт поддержки протокола не задаёт методику отбора источников, контроля качества контекста и доказательной проверки результата.

Современные исследования MCP подчёркивают его роль как унифицированного интерфейса между моделями, инструментами и ресурсами, а также указывают на проблемы доверия, безопасности и разграничения полномочий при развёртывании MCP-серверов [3, 4]. Близкие направления развиваются в работах по retrieval-augmented generation, агентным системам и инструментально-расширенным языковым моделям [5–8]. Вопросы доверия к системам искусственного интеллекта и интеграции больших языковых моделей в контуры анализа разнородных данных активно обсуждаются и в русскоязычной литературе [9, 10]. В маркетинговой аналитике большие языковые модели применяются для дополнения данных опросов, сегментации потребителей, анализа тональности и объяснения предпочтений [11–13]. Тем не менее, задача построения воспроизводимого MCP-контура маркетингового исследования, где каждый результат связан с источником и может быть повторно проверен, остаётся недостаточно формализованной.

Цель данной работы заключается в разработке и экспериментальной проверке метода доверенной контекстуализации MCP-Trust для интеллектуальных агентов маркетинговой аналитики. Объектом исследования является процесс формирования аналитического ответа агентом, подключенным к

нескольким МСР-серверам. Предметом исследования является метод выбора и согласования контекстных данных с учётом релевантности, свежести, надёжности, стоимости и доказательной прослеживаемости источников.

Научная новизна работы состоит в следующем. Во-первых, предложено понятие контекстного пакета МСР-Trust, объединяющего запрос, схему результата, полномочия пользователя, бюджеты контекста, реестр серверов, доказательства и трассировку вычислений. Во-вторых, предложена функция ранжирования МСР-серверов, учитывающая релевантность, свежесть, надёжность, доступность и стоимость вызова. В-третьих, введена процедура согласования противоречивых источников через показатель согласованности и доверия к результату. В-четвертых, предложена интегральная метрика эффективности, одновременно учитывающая точность, покрытие источниками, согласованность, воспроизводимость и стоимость вычисления.

I. ПОСТАНОВКА ЗАДАЧИ

Рассмотрим интеллектуального агента маркетинговой аналитики, работающего в среде МСР. В общем случае такая среда может быть представлена ориентированным графом:

$$G = \{H, C, S, T, D, E\}, \quad (1)$$

где H – хост-приложение агента; C – множество МСР-клиентов внутри хоста; S – множество МСР-серверов; T – множество инструментов, предоставляемых серверами; D – источники данных; E – доказательства происхождения данных и связей между ними.

Вершины графа соответствуют приложениям, серверам, инструментам и источникам, а рёбра описывают допустимые вызовы и передачу контекстной информации.

Для отдельного пользовательского запроса q агент формирует контекстный пакет:

$$B = \langle q, x, \sigma, u, L, A, E, \pi \rangle, \quad (2)$$

где смысл и назначение элементов пакета B раскрыты в Табл. I.

ТАБЛИЦА I
ЭЛЕМЕНТЫ КОНТЕКСТНОГО ПАКЕТА МСР-TRUST

Символ	Содержание
q	аналитический запрос пользователя
x	частично нормализованные входные данные
σ	схема результата и допустимые типы полей
u	профиль полномочий и ограничений пользователя
L	бюджеты времени, стоимости и объёма контекста
A	реестр доступных МСР-серверов и инструментов
E	множество доказательств и ссылок на источники
π	журнал трассировки для повторного воспроизведения

Вызов i -го инструмента или сервера в рамках пакета B задаётся отображением:

$$T_i(B) \rightarrow y_i = \langle z_i, e_i, m_i, c_i, \tau_i \rangle, \quad (3)$$

где z_i – значение одного и того же аналитического поля, извлечённое или вычисленное i -м сервером; e_i – ссылка на источник либо фрагмент доказательства; m_i – метаданные результата; c_i – стоимость обращения; τ_i – задержка.

Требуется выбрать подмножество серверов $S^* \subseteq S$, где S – множество МСР-серверов среды (1), и порядок их вызова так, чтобы итоговый аналитический результат, агрегирующий извлечённые значения z_i из (3), обладал минимальной ошибкой относительно проверяемых данных и максимальной доверенной прослеживаемостью при соблюдении ограничений L , входящих в контекстный пакет B (2). Формальная постановка выбора S^* приведена ниже в (5).

Прикладная постановка задачи ориентирована на типовые действия маркетингового аналитика: оценку размера сегмента, сравнение каналов продаж, выявление отклонений в динамике спроса, формирование профиля целевой аудитории и подготовку таблиц с числовыми показателями. В таких задачах агент должен не только сгенерировать текстовое объяснение, но и указать источники, на основании которых получены отдельные значения.

II. ТЕОРИЯ

A. Ранжирование МСР-серверов

Пусть для каждого МСР-сервера s_i известны или оцениваются частные показатели: релевантность $rel_i(q)$ к запросу, свежесть $fresh_i(t)$ относительно момента t , надёжность $\rho_i(t)$, доступность $avail_i(t)$ и нормированная стоимость $cost_i$. Тогда приоритет обращения к серверу определяется выражением:

$$R_i(q, t) = \alpha rel_i(q) + \beta fresh_i(t) + \gamma \rho_i(t) + \delta avail_i(t) - \eta cost_i, \quad (4)$$

где $\alpha, \beta, \gamma, \delta, \eta$ – неотрицательные весовые коэффициенты, задаваемые политикой агентной системы. Для маркетинговых данных коэффициент β возрастает при анализе быстро меняющихся рынков, тогда как коэффициент γ должен преобладать при формировании отчётов, используемых для управленческих решений.

Выбор серверов формулируется как дискретная оптимизационная задача:

$$S^* = \arg \max \sum R_i(q, t), \text{ при } \sum c_i \leq C_{max}, \sum \tau_i \leq T_{max}. \quad (5)$$

Здесь C_{max} и T_{max} – ограничения на стоимость и задержку, заданные вектором ограничений L контекстного пакета B (2). В разработанном методе используется жадная эвристика: сначала выбираются обязательные источники по схеме результата σ , затем добавляются резервные источники с наибольшим R_i до исчерпания бюджета.

В. Согласование источников и доказательная трассировка

Для набора результатов $Y = \{y_1, \dots, y_m\}$, относящихся к одному аналитическому полю, вводится показатель согласованности:

$$Cons(Y) = 1 - \min(1, \text{stdev}(z_1, \dots, z_m) / (|\text{mean}(z_1, \dots, z_m)| + \varepsilon)), \quad (6)$$

где $z_1, \dots, z_m$ – значения одного и того же аналитического поля, извлечённые или вычисленные разными серверами согласно (3); $\text{mean}(z_1, \dots, z_m)$ и $\text{stdev}(z_1, \dots, z_m)$ – соответственно среднее значение и стандартное отклонение; ε – малая положительная константа ($\varepsilon \ll 1$; в экспериментах $\varepsilon = 10^{-6}$), добавляемая в знаменатель для защиты от деления на ноль при среднем значении, близком к нулю.

Если $Cons(Y)$ ниже порогового значения, агент не обязан выбирать единственное число: он формирует предупреждение о конфликте источников и включает в отчёт диапазон значений с указанием причин расхождения. Это отличает доверенную контекстуализацию от обычного сценария, где языковая модель может усреднить противоречивые значения без объяснения происхождения данных.

Итоговое доверие к аналитическому значению z оценивается как:

$$\text{Trust}(z) = \min\{1, Cons(Y) \cdot Pe \cdot [\sum_{i=1}^m (w_i \cdot \rho_i)] / [\sum_{i=1}^m (w_i) + \varepsilon]\} \quad (7)$$

где w_i – вес i -го источника ($i = 1, \dots, m$); ρ_i – надёжность i -го источника, то есть значение введённой ранее функции $\rho_i(t)$ на момент формирования ответа (аргумент t далее опущен для краткости); $Cons(Y)$ и Pe – общие для оцениваемого значения множители; ε – та же малая положительная константа, что и в (6).

Для хранения доказательств используется ориентированный граф провенанса, совместимый по смыслу с моделью PROV-O [14]: вершины описывают источники, вызовы инструментов, промежуточные таблицы и итоговые поля отчёта, а рёбра связывают результат с порождающими его операциями.

С. Интегральная метрика эффективности

Для сравнения вариантов агентной архитектуры предложена интегральная метрика:

$$F = Acc \cdot Pe \cdot Cons \cdot Replay / (1 + Cnorm), \quad (8)$$

где $Acc = 1 - MAPE/100$ – нормированная точность; Pe – покрытие полей результата доказательствами; $Cons$ – средняя согласованность источников; $Replay$ – доля успешно воспроизведённых результатов при повторном запуске; $Cnorm$ – нормированная стоимость обращений к внешним инструментам. Метрика F возрастает только в том случае, если повышение точности не достигается ценой полной потери воспроизводимости или чрезмерного роста стоимости. Обобщенная архитектура метода MCP-Trust показана на Рис. 1.

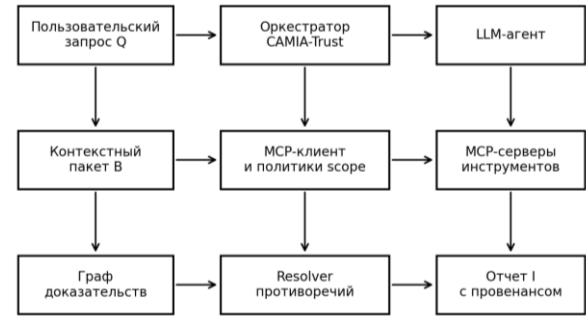


Рис. 1. Обобщенная архитектура метода MCP-Trust

Сравнение исследуемых режимов по точности и по интегральной метрике эффективности приведено на Рис. 2 и Рис. 3.

III. РЕЗУЛЬТАТЫ ЭКСПЕРИМЕНТОВ

Для проверки предложенного метода создан имитационный стенд на языке Python 3.11 с использованием библиотек numpy, pandas и matplotlib. Стенд моделирует работу агента, получающего маркетинговые запросы и обращающегося к виртуальным MCP-серверам. Каждый сервер характеризуется тематикой, надёжностью, средней задержкой, стоимостью вызова, вероятностью отказа и вероятностью предоставления доказательства.

Синтетические данные формировались программным генератором с фиксированным зерном генерации ($seed = 20260530$, Табл. II), что обеспечивает полную воспроизводимость стенда. Для каждого из 180 запросов случайным образом выбиралась одна из шести отраслевых категорий и разыгрывалось эталонное значение целевого показателя, относительно которого затем вычислялась ошибка ответа.

Ответ каждого виртуального MCP-сервера моделировался как эталонное значение с мультипликативным шумом, дисперсия которого обратно пропорциональна надёжности сервера, а устаревшие источники дополнительно получали систематическое смещение.

Отказ сервера и наличие доказательства происхождения разыгрывались как случайные события с заданными вероятностями; стоимость и задержка вызова выбирались из фиксированных для каждого сервера диапазонов. Итоговые метрики усреднялись по всем 180 запросам.

Эксперимент не претендует на замену промышленной апробации на реальных закрытых данных компаний. Его назначение – воспроизводимо проверить, как изменение стратегии отбора и согласования источников влияет на точность, доказательность и стоимость результата при контролируемых параметрах среды.

Параметры имитационного стенда приведены в Табл. II.

ТАБЛИЦА II
ПАРАМЕТРЫ ЭКСПЕРИМЕНТАЛЬНОГО СТЕНДА

Параметр	Символ	Значение
Количество запросов	Nq	180
Количество виртуальных MCP-серверов	Ns	12
Число отраслевых категорий	Kc	6
Случайное зерно генерации	$seed$	20260530
Сравниваемые режимы	-	LLM без MCP; MCP-static; CAMIA; MCP-Trust
Метрики	-	MAPE, MAE, $Cons$, Pe , $Cnorm$, τ , $Replay$, F

Сводные результаты по сравниваемым режимам – в Табл. III.

Сравнивались четыре режима: 1) LLM без MCP, где результат формируется без структурированного обращения к источникам; 2) MCP-static, где используется фиксированный набор серверов без доверенного ранжирования; 3) CAMIA, соответствующий адаптивной контекстуализации из исходного варианта статьи; 4) MCP-Trust, дополняющий CAMIA ранжированием надёжности, доказательным графом, проверкой конфликтов и трассировкой воспроизведения. Режимы сравнивались по полному набору метрик эксперимента, перечисленных в Табл. II: ошибкам MAPE (средняя абсолютная процентная ошибка) и MAE (средняя абсолютная ошибка в единицах показателя), согласованности $Cons$, покрытию доказательствами Pe , нормированной стоимости $Cnorm$, средней задержке τ , воспроизводимости $Replay$ и интегральной метрике F ; значения всех метрик приведены в Табл. III.

ТАБЛИЦА III
РЕЗУЛЬТАТЫ ИМИТАЦИОННОГО ЭКСПЕРИМЕНТА

Метод	MAPE, %	MAE, у.е.	$Cons$	Pe	$Cnorm$	τ , с	$Replay$	F
LLM без MCP	15,93	14,62	0,350	0,233	0,168	0,82	0,46	0,027
MCP-static	9,76	9,05	0,446	0,915	0,551	1,94	0,80	0,190
CAMIA	8,83	8,17	0,464	0,965	0,629	2,57	0,91	0,228
MCP-Trust	8,28	7,64	0,554	0,971	0,809	3,41	0,99	0,270

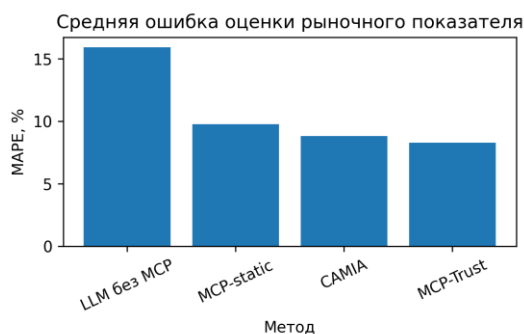


Рис. 2. Средняя абсолютная процентная ошибка в сравниваемых режимах

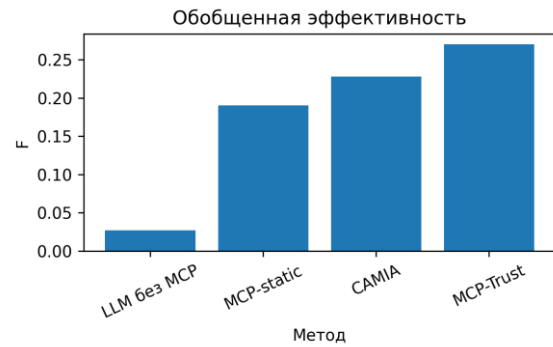


Рис. 3. Интегральная метрика эффективности в сравниваемых режимах

Полученные результаты показывают последовательное улучшение качества при переходе от неструктурированного взаимодействия к управляемой MCP-архитектуре. В режиме без MCP средняя ошибка MAPE составила 15,93 %. Статическое подключение MCP-серверов снизило ошибку до 9,76 %, CAMIA – до 8,83 %, а MCP-Trust – до 8,28 %. Таким образом, доверенная контекстуализация уменьшила MAPE на 48,0 % относительно режима без MCP и на 6,2 % относительно CAMIA.

Рост точности сопровождается улучшением доказательной прослеживаемости. Покрытие источниками Pe увеличилось с 0,233 в базовом режиме до 0,971 в MCP-Trust. Коэффициент воспроизводимости $Replay$ вырос с 0,46 до 0,99. Средняя согласованность источников увеличилась до 0,554, что выше показателей MCP-static и CAMIA. При этом стоимость и задержка MCP-Trust ожидаемо выше ($Cnorm = 0,809$ против 0,168 у режима без MCP, средняя задержка $\tau = 3,41$ с против 0,82 с) – метод выполняет дополнительные проверки, обращается к резервным источникам и сохраняет доказательный граф.

IV. ОБСУЖДЕНИЕ РЕЗУЛЬТАТОВ

Положительный эффект MCP-Trust объясняется тем, что метод отделяет протокольный уровень подключения инструментов от методического уровня формирования доверенного контекста. MCP обеспечивает унифицированный канал обмена с серверами, однако без правил выбора источников и проверки доказательств агент остаётся уязвимым к устаревшим, нерелевантным или противоречивым данным. Предложенный контекстный пакет B делает процесс получения результата более управляемым: для каждого значения фиксируются источник, параметры вызова и связь с итоговым отчётом.

Сравнение с CAMIA показывает, что адаптивная контекстуализация сама по себе повышает качество результата, однако добавление доверенного ранжирования и трассировки дополнительно увеличивает интегральную метрику F с 0,228 до 0,270, то есть примерно на 18,4 %. По сравнению со

статической МСР-схемой прирост F составил около 42,1 %. Это означает, что основной вклад метода связан не только с большим числом источников, но и с управлением их надёжностью и воспроизводимостью.

Практическое значение результатов заключается в возможности построения интеллектуального помощника аналитика, который формирует не просто текстовый вывод, а проверяемый аналитический артефакт: таблицу, график, резюме и список источников для каждой группы значений. Такой подход полезен при подготовке отчётов о целевой аудитории, динамике спроса, конкурентной среде и эффективности каналов коммуникации.

Ограничением работы является имитационный характер эксперимента. В синтетическом стенде известны распределения ошибок и надёжности серверов, тогда как в промышленной среде эти параметры должны оцениваться по истории обращений, экспертным рейтингам и аудитам данных. Кроме того, в статье не моделировались целенаправленные атаки на инструменты, включая подмену описаний, скрытую эксфильтрацию данных и конфликт полномочий пользователя. Эти вопросы требуют отдельного исследования с учётом современных угроз МСР-среды.

ВЫВОДЫ И ЗАКЛЮЧЕНИЕ

1. Предложен метод МСР-Trust, расширяющий применение Model Context Protocol в интеллектуальных агентах маркетинговой аналитики за счёт формализованного контекстного пакета, ранжирования серверов, доказательного графа и трассировки повторного воспроизведения.

2. Разработан математический аппарат выбора МСР-серверов и оценки доверия к аналитическому результату. Введены функции $R_i(q, t)$, $Cons(Y)$, $Trust(z)$ и интегральная метрика F , позволяющая сравнивать агентные архитектуры не только по точности, но и по доказательности, согласованности, воспроизводимости и стоимости.

3. Имитационный эксперимент показал, что МСР-Trust снижает МАРЕ до 8,28 % против 15,93 % у режима без МСР, повышает покрытие источниками до 0,971 и обеспечивает коэффициент воспроизводимости 0,99. Интегральная метрика F увеличилась до 0,270, что выше результатов SAMIA и статической МСР-схемы.

4. Дальнейшее развитие работы целесообразно связать с проверкой метода на реальных корпоративных наборах маркетинговых данных, с анализом устойчивости к вредоносным МСР-серверам и с разработкой программного модуля аудита контекстных пакетов для промышленных аналитических платформ.

ЛИТЕРАТУРА

[1] Anthropic. Introducing the Model Context Protocol [Electronic resource]. 2024. URL: <https://www.anthropic.com/news/model-context-protocol> (дата обращения: 30.05.2026).

- [2] Model Context Protocol. Specification [Electronic resource]. 2025. URL: <https://modelcontextprotocol.io/specification/2025-03-26> (дата обращения: 30.05.2026).
- [3] Hou X., Zhao Y., Wang S., Wang H. Model Context Protocol (MCP): Landscape, Security Threats, and Future Research Directions. arXiv:2503.23278. 2025. DOI: 10.48550/arXiv.2503.23278.
- [4] Narajala V.S., Habler I. Enterprise-Grade Security for the Model Context Protocol (MCP): Frameworks and Mitigation Strategies. arXiv:2504.08623. 2025. DOI: 10.48550/arXiv.2504.08623.
- [5] Lewis P., Perez E., Piktus A., Petroni F., Karpukhin V., Goyal N. et al. Retrieval-Augmented Generation for Knowledge-Intensive NLP Tasks // Advances in Neural Information Processing Systems. 2020. Vol. 33. P. 9459–9474.
- [6] Yao S., Zhao J., Yu D., Du N., Shafran I., Narasimhan K., Cao Y. ReAct: Synergizing Reasoning and Acting in Language Models // Proc. ICLR. 2023.
- [7] Xi Z. et al. The rise and potential of large language model based agents: a survey // Science China Information Sciences. 2025. Vol. 68. Art. 121101. DOI: 10.1007/s11432-024-4222-0.
- [8] Shen Z. LLM With Tools: A Survey. arXiv:2409.18807. 2024. DOI: 10.48550/arXiv.2409.18807.
- [9] Гарбук С.В. Особенности применения понятия «доверие» в области искусственного интеллекта // Искусственный интеллект и принятие решений. 2020. № 3. С. 15–21.
- [10] Федоров А.М., Датьев И.О., Вишняков И.Г. Метод интеграции больших языковых моделей в алгоритмы фокусированного мониторинга открытых данных социальных медиа // Информатика и автоматизация. 2025. Т. 24, № 6. С. 1623–1648. DOI: 10.15622/ia.24.6.4.
- [11] Wang M., Zhang D.J., Zhang H. Large Language Models for Market Research: A Data-augmentation Approach. arXiv:2412.19363. 2026. DOI: 10.48550/arXiv.2412.19363.
- [12] Li Y., Liu Y., Yu M. Consumer segmentation with large language models // Journal of Retailing and Consumer Services. 2025. Vol. 82. Art. 104078. DOI: 10.1016/j.jretconser.2024.104078.
- [13] Krugmann J.O., Hartmann J. Sentiment Analysis in the Age of Generative AI // Customer Needs and Solutions. 2024. Vol. 11. Art. 3. DOI: 10.1007/s40547-024-00143-4.
- [14] Lebo T., Sahoo S., McGuinness D. et al. PROV-O: The PROV Ontology. W3C Recommendation. 2013. URL: <https://www.w3.org/TR/prov-o/> (дата обращения: 30.05.2026).

Информация об авторах

Нижниченко Георгий Алексеевич, аспирант-выпускник, Московский финансово-юридический университет МФЮА, Москва, Россия, e-mail: nizhge@gmail.com.

Павлов Валерий Анатольевич, кандидат экономических наук, доцент кафедры информационных систем, технологий и автоматизации в строительстве НИУ МГСУ, Москва, Россия, e-mail: 29359332@mfua.ru.

Trust-calibrated Model Context Protocol (MCP) contextualization method for intelligent marketing analytics agents

Georgy Nizhnichenko¹, Valery Pavlov²

¹ Moscow University of Finance and Law MFUA, Moscow, Russia

² Moscow State University of Civil Engineering, Moscow, Russia

Abstract – The paper proposes the MCP-Trust method for trust-calibrated contextualization of intelligent marketing

analytics agents. The method extends the use of the Model Context Protocol by introducing a formal context package, MCP server ranking, source consistency checking, evidence graph construction and replay tracing. A mathematical apparatus is proposed for selecting MCP servers and estimating the trust of analytical outputs. A simulation experiment with 180 synthetic analytical requests and 12 virtual MCP servers showed that MCP-Trust reduced MAPE from 15.93 to 8.28%, increased evidence coverage from 0.233 to 0.971 and increased replay reproducibility from 0.46 to 0.99. The integral efficiency metric reached 0.270 compared with 0.228 for CAMIA and 0.190 for the static MCP scheme. The results confirm the applicability of trust-calibrated contextualization for reproducible market intelligence systems.

Keywords – Model Context Protocol, intelligent agents, marketing analytics, evidence verification, data provenance.

REFERENCES

- [1] Anthropic. Introducing the Model Context Protocol [Electronic resource]. 2024. Available at: <https://www.anthropic.com/news/model-context-protocol> (accessed: 30.05.2026).
- [2] Model Context Protocol. Specification [Electronic resource]. 2025. Available at: <https://modelcontextprotocol.io/specification/2025-03-26> (accessed: 30.05.2026).
- [3] Hou X., Zhao Y., Wang S., Wang H. Model Context Protocol (MCP): Landscape, Security Threats, and Future Research Directions. arXiv:2503.23278. 2025. DOI: 10.48550/arXiv.2503.23278.
- [4] Narajala V.S., Habler I. Enterprise-Grade Security for the Model Context Protocol (MCP): Frameworks and Mitigation Strategies. arXiv:2504.08623. 2025. DOI: 10.48550/arXiv.2504.08623.
- [5] Lewis P., Perez E., Piktus A., Petroni F., Karpukhin V., Goyal N. et al. Retrieval-Augmented Generation for Knowledge-Intensive NLP Tasks. Advances in Neural Information Processing Systems. 2020. Vol. 33. P. 9459–9474.
- [6] Yao S., Zhao J., Yu D., Du N., Shafran I., Narasimhan K., Cao Y. ReAct: Synergizing Reasoning and Acting in Language Models. Proc. ICLR. 2023.
- [7] Xi Z. et al. The rise and potential of large language model based agents: a survey. Science China Information Sciences. 2025. Vol. 68. Art. 121101. DOI: 10.1007/s11432-024-4222-0.
- [8] Shen Z. LLM With Tools: A Survey. arXiv:2409.18807. 2024. DOI: 10.48550/arXiv.2409.18807.
- [9] Garbuk S.V. Osobennosti primeneniya ponyatiya «doverie» v oblasti iskusstvennogo intellekta [Features of applying the concept of “trust” in the field of artificial intelligence]. Iskuststvennyi intellekt i prinyatie reshenii [Artificial Intelligence and Decision Making]. 2020. No. 3. P. 15–21. (In Russian).
- [10] Fedorov A.M., Dat'ev I.O., Vishnyakov I.G. Metod integracii bol'shih yazykovykh modelej v algoritmy fokusirovannogo monitoringa otkrytykh dannyh social'nyh media // Informatika i avtomatizaciya. 2025. T. 24, № 6. S. 1623–1648. DOI: 10.15622/ia.24.6.4. (In Russian).
- [11] Wang M., Zhang D.J., Zhang H. Large Language Models for Market Research: A Data-augmentation Approach. arXiv:2412.19363. 2026. DOI: 10.48550/arXiv.2412.19363.
- [12] Li Y., Liu Y., Yu M. Consumer segmentation with large language models. Journal of Retailing and Consumer Services. 2025. Vol. 82. Art. 104078. DOI: 10.1016/j.jretconser.2024.104078.
- [13] Krugmann J.O., Hartmann J. Sentiment Analysis in the Age of Generative AI. Customer Needs and Solutions. 2024. Vol. 11. Art. 3. DOI: 10.1007/s40547-024-00143-4.
- [14] Lebo T., Sahoo S., McGuinness D. et al. PROV-O: The PROV Ontology. W3C Recommendation. 2013. Available at: <https://www.w3.org/TR/prov-o/> (accessed: 30.05.2026).

Information about the authors

Georgy A. Nizhnichenko, postgraduate student, Moscow University of Finance and Law MFUA, Moscow, Russia, e-mail: nizhge@gmail.com.

Valery A. Pavlov, Candidate of Economic Sciences, Associate Professor of the Department of Information Systems, Technologies and Automation in Construction, Moscow State University of Civil Engineering, Moscow, Russia, e-mail: 29359332@s.mfua.ru.