

УДК 004.89:004.7

Применение больших языковых моделей (LLM) для диагностики сетевых инцидентов

Е.А. Колесников✉, В.Н. Кострова

Воронежский институт высоких технологий, Воронеж, Россия

В статье рассматривается применение больших языковых моделей для диагностики сетевых инцидентов в рамках AIOps. Показано, что традиционные методы управления сетями становятся неэффективными в условиях экспоненциального роста телеметрических данных. Анализируются ключевые направления использования LLM: автоматизированный анализ журналов, локализация корневых причин отказов, взаимодействие с сетью на естественном языке и мультиагентные системы. Рассматриваются ограничения технологии: галлюцинации LLM, чрезмерная агентность, уязвимость к атакам, вопросы конфиденциальности и вычислительная сложность. Сделан вывод, что LLM открывают новые возможности для интеллектуальной автоматизации сетевой диагностики.

Ключевые слова: большие языковые модели, LLM, AIOps, диагностика сетевых инцидентов, анализ журналов, локализация корневых причин, агентный искусственный интеллект.

Application of Large Language Models (LLMs) for Network Incident Diagnosis

Е.А. Kolesnikov✉, V.N. Kostrova

Voronezh Institute of High Technologies, Voronezh, Russia

The article examines the application of large language models for network incident diagnosis within AIOps. Traditional network management methods are shown to become ineffective under exponentially growing telemetry data volumes. Key LLM applications are analyzed: automated log analysis, localization of root causes of failures, natural language network interaction, and multi-agent systems. Technology limitations are considered: LLM hallucinations, excessive agency, vulnerability to attacks, privacy concerns, and computational complexity. It is concluded that LLMs open new opportunities for intelligent automation of network diagnostics.

Keywords: large language models, LLM, AIOps, network incident diagnosis, log analysis, root cause localization, agentic artificial intelligence.

Введение

Современные сетевые инфраструктуры характеризуются беспрецедентным уровнем сложности. Переход к распределённым архитектурам, внедрение облачных технологий, рост числа подключённых устройств и появление новых стандартов связи привели к экспоненциальному увеличению объёмов генерируемых операционных данных [1]. Традиционные подходы к мониторингу и диагностике инцидентов, основанные на ручном анализе журналов событий и статических правилах, перестали быть эффективными. Инженеры сталкиваются с лавинообразным потоком алертов – в некоторых случаях IT-команды получают тысячи уведомлений в час [2].

В этих условиях на первый план выходит концепция AIOps (Artificial Intelligence for IT Operations) – применение технологий искусственного интеллекта для автоматизации и повышения качества операционного управления

ИТ-инфраструктурами [3]. На протяжении последних лет машинное обучение было ключевым компонентом инструментов сетевого управления. Однако появление больших языковых моделей (LLM) добавило принципиально новые возможности на пути к сетевым автономиям [4].

Целью настоящей работы является систематизация современных представлений о применении больших языковых моделей для диагностики сетевых инцидентов, анализ ключевых возможностей, предоставляемых LLM в рамках AIOps, и выявление существующих ограничений, препятствующих их повсеместному внедрению.

Современное состояние и предпосылки внедрения LLM в сетевую диагностику

Языковые модели в настоящее время всё активнее интегрируются в современные AIOps-решения, особенно в форме LLM-агентов, которые способны интерпретировать телеметрические данные и предпринимать корректирующие действия с минимальным вмешательством человека [5]. В отличие от традиционных подходов, LLM демонстрируют способность к семантическому пониманию и генерации естественного языка, что открывает возможности для качественно нового уровня взаимодействия человека с сетью.

Ключевым драйвером внедрения LLM является критический уровень сложности современных сетей, при котором традиционные системы мониторинга генерируют огромное количество алертов, часто ложных, что приводит к «усталости от оповещений» и снижению эффективности работы инженеров. При этом сетевая диагностика по-прежнему остаётся трудоёмкой задачей, требующей ручного анализа топологических и временных взаимосвязей между событиями и устройствами [6].

Значимым шагом вперёд стало появление специализированных бенчмарков для оценки LLM-агентов в задачах сетевой диагностики. В работе [7] представлен бенчмарк NIKA, включающий сотни сетевых инцидентов, охватывающих пять сетевых сценариев – от центров обработки данных до сетей интернет-провайдеров, и покрывающий 54 репрезентативных типа сетевых проблем [7].

Ключевые возможности применения LLM для диагностики сетевых инцидентов

Анализ современных решений и публикаций позволяет выделить несколько ключевых направлений, в которых LLM демонстрируют наибольшую эффективность в задачах сетевой диагностики.

Анализ и интерпретация журналов событий. Сетевые журналы являются одним из основных источников данных для диагностики инцидентов. Однако их объёмы, разнообразие форматов и семантика делают ручной анализ крайне трудоёмким. Крупные облачные системы генерируют 30–50 ГБ журналов в час, что делает применение LLM для каждого сообщения экономически неоправданным [8].

Для решения этой проблемы предложен фреймворк **LogRules**, который использует генерацию правил на основе LLM. LogRules состоит из трёх этапов: индукции (создание правил с помощью GPT-4o-mini), выравнивания (дообучение компактных LLM, например LLaMA-3-8B, через контрастную оптимизацию предпочтений) и вывода (анализ с использованием репозитория правил). Эксперименты показали, что LogRules превосходит существующие LLM-методы в задачах парсинга логов и обнаружения аномалий, а также демонстрирует бóльшую устойчивость по сравнению с подходами на примерах [8]. Ключевое преимущество – снижение стоимости анализа примерно на 88% за счёт использования меньших моделей [8].

Для промышленных сетевых платформ разработана система **RT-LogAAS** (Real-time Log Anomaly Analysis System). Она объединяет потоковый движок реального времени на базе Kafka и Flink с технологией LogLSD, использующей Locality-Sensitive Hashing для агрегации логов, и LLM для семантической интерпретации. Такой гибридный подход эффективно обрабатывает гетерогенные логи от разных производителей и типов устройств, демонстрируя высокую применимость в производственных средах [9].

Локализация корневых причин отказов. Локализация корневых причин отказов. В работе [10] предложена система ALPHA, представляющая собой конвейер активного обучения для полностью автоматизированного анализа логов. ALPHA объединяет семантическое встраивание, кластеризацию с репрезентативной выборкой и автоматическую аннотацию с помощью LLM. Двухэтапная стратегия дообучения с адаптивным выбором промптов на основе наблюдаемых ошибок LLM позволяет достичь точности обнаружения аномалий, сопоставимой с методами полного контроля, при практически полном исключении участия человека в процессе разметки [10].

Система **MicroRCA-Agent** предлагает подход к анализу первопричин отказов в микросервисных архитектурах на основе LLM-агентов с multimodal data fusion. Решение комбинирует алгоритм парсинга логов Drain с многоуровневой фильтрацией для сжатия логов, метод Isolation Forest для обнаружения аномалий в трейсах и двухэтапную стратегию анализа метрик с фильтрацией на основе статистического коэффициента симметрии. Модальный модуль анализа использует специально разработанные кросс-модальные промпты для глубокой интеграции разномодальной информации. Валидация на реальных сценариях отказов подтвердила эффективность подхода [11].

Взаимодействие с сетью на естественном языке. Наиболее заметным с практической точки зрения прорывом стало внедрение LLM-интерфейсов, позволяющих инженерам взаимодействовать с сетью на естественном языке, без необходимости изучения сложных синтаксисов командных строк или языков запросов. Крупнейшие вендоры в области сетевого оборудования и программного обеспечения уже интегрировали подобные решения в свои продукты.

Cisco AI Assistant использует предварительно обученные модели, настроенные для специфических сетевых задач. Ассистент обеспечивает управляемое устранение неполадок, упрощает анализ корневых причин и ускоряет решение проблем. Он интегрирован с платформами Meraki и ThousandEyes, предоставляя унифицированный доступ к данным о производительности и состоянии сети. Cisco AI Assistant способен обрабатывать сложные запросы, анализируя информацию о сети, приложениях, путях передачи данных и обнаруженных событиях для точной локализации проблемного домена [2, 12].

HPE Juniper развивает своего виртуального ассистента Marvis, который использует обработку естественного языка для взаимодействия с сетевыми операторами. Благодаря обновлённой агентной архитектуре Marvis обеспечивает устранение неполадок в реальном времени, а самоуправляемые агенты сотрудничают в различных сетевых доменах и приложениях. Marvis интегрирует несколько AI-агентов в свой интерфейс, что позволяет вести более естественные, контекстные диалоги и предоставлять точные ответы, заменяя традиционные информационные панели и интерфейсы командной строки интуитивно понятными, ориентированными на задачи разговорами [13, 14].

Nokia внедрила в свою платформу EDA (Event-Driven Architecture) интерфейс Ask EDA, который часто сравнивают с «ChatGPT для сети». Ассистент позволяет

операторам задавать вопросы на простом английском языке, например: «перечислить мои SRL-интерфейсы» или «включена ли BFD?», и получать в ответ живые структурированные данные, таблицы и даже визуализации. Ключевой особенностью подхода Nokia является гибкая агентная структура, где центральный LLM выступает в роли «мозга», координирующего набор подключаемых агентов – специализированных инструментов. Это позволяет добавлять новые возможности и рабочие процессы без переобучения модели. Такой подход решает проблему галлюцинаций, так как ответы модели основаны на фактических данных из живой телеметрии и инструментов [15].

Мультиагентные системы для самовосстановления сетей. TM Forum в проекте **Project ONE** развивает концепцию агентной ODA (Open Digital Architecture), где AI-агенты способны вести переговоры с клиентами о компенсациях, взаимодействуя через Model Context Protocol (MCP) и Agent-to-Agent Protocol (A2A) с компонентами ODA. Ключевая инновация – использование Guardrails (ограничителей), например, разрешение агенту компенсировать клиента только до определённой суммы без одобрения человека [16].

China Mobile достигла уровня 4 автономных сетей по классификации TM Forum в своих центрах сетевых операций. Компания разработала два типа AI-агентов (ролевые copilot-ы и сценарные агенты) и построила большую модель с более чем 10 миллиардами параметров, а также замкнутые автономные инструментальные цепочки на основе коллаборации больших и малых моделей. Система интегрирована с технологией RAG, что позволило повысить точность ответов до 90% и сократить галлюцинации до 3%. Результаты внедрения впечатляют: сокращение среднего времени восстановления на 30%, повышение качества обслуживания, а также высвобождение эквивалента 5 500 полных рабочих позиций (75% в сетевой эксплуатации, 25% в работе с клиентами). В провинции Чжэцзян внедрение привело к сокращению MTTR по отказам RAN на 27% и на 87% по отказам базовой сети [16].

Ограничения и вызовы при внедрении LLM

Несмотря на впечатляющие возможности, внедрение LLM для диагностики сетевых инцидентов сопряжено с рядом существенных ограничений и вызовов, которые необходимо учитывать.

Галлюцинации и достоверность решений. Одной из ключевых проблем при использовании LLM в высокоответственных сетевых средах является склонность к галлюцинациям – генерации правдоподобных, но фактически неверных или неподтверждённых результатов. В работе «When AIOps Become 'AI Oops'» было показано, что современные AIOps-агенты могут принимать ошибочные решения, будучи введёнными в заблуждение сгенерированными иллюзиями [4]. Исследователи отмечают, что оценка LLM-агентов на бенчмарке NIKA показала: более крупные модели чаще успешно обнаруживают проблемы, но всё ещё испытывают трудности с локализацией отказов и определением корневых причин [7].

Основной подход к минимизации галлюцинаций при сетевой диагностике заключается в grounding – обеспечении того, чтобы все ответы модели основывались на фактических данных из внешних источников, таких как конфигурационные файлы, журналы событий и метрики телеметрии. Такой подход, используемый в Nokia Ask EDA и других системах, позволяет использовать LLM для композиции ответов, но извлекает фактические данные из проверяемых источников [15].

Чрезмерная агентность и вопросы безопасности. При предоставлении LLM-агентам возможности автоматически выполнять действия в сети возникает риск чрезмерной агентности – предоставления системе больше прав, полномочий или

автономии, чем необходимо для выполнения поставленных задач. Согласно спецификации OWASP LLM06:2025, чрезмерная агентность является одной из критических уязвимостей LLM-приложений. Она возникает, когда система в ответ на неожиданные, двусмысленные или манипулируемые входные данные может выполнить вредоносные действия [17].

Тремя корневыми причинами чрезмерной агентности являются: избыточная функциональность (доступ к ненужным инструментам), избыточные разрешения (слишком широкие права доступа к системам) и избыточная автономия (выполнение опасных операций без подтверждения человеком). Исследование RSA Conference продемонстрировало уязвимость систем AIOps к манипуляции телеметрией, когда злоумышленники могут внедрять специально сформированные данные, чтобы вводить AI-агентов в заблуждение и заставлять их предпринимать действия, нарушающие целостность инфраструктуры. Разработанная авторами атака AIOpsDoom полностью автоматизирована – она сочетает разведку, фаззинг и генерацию состязательных входных данных с помощью LLM и может работать без предварительного знания целевой системы [4].

Конфиденциальность и безопасность данных. Использование LLM для анализа журналов событий поднимает серьёзные вопросы защиты конфиденциальной информации. Сетевые журналы часто содержат чувствительные данные о пользователях, бизнес-процессах и внутренней архитектуре. При передаче таких данных в облачные LLM-сервисы или даже при локальном развёртывании существует риск утечки. Исследователи из Мельбурнского университета отмечают, что применение лёгких LLM, развёрнутых локально, может значительно сократить время реагирования на инциденты и снизить риски утечки данных [18]. Дообучение компактных моделей на внутренних данных организации позволяет совместить эффективность LLM с требованиями безопасности.

Вычислительная сложность и стоимость развёртывания. Применение LLM в режиме реального времени для анализа больших объёмов сетевых данных требует значительных вычислительных ресурсов. Запрос к API крупной языковой модели имеет стоимость, а при масштабах крупного сетевого оператора затраты могут стать существенными. Хотя современные системы используют агентную архитектуру, в которой специализированные модели (иногда меньшего размера) выполняют конкретные задачи под управлением центрального LLM, проблема затрат остаётся открытой. Разработка эффективных методов дообучения и сжатия моделей, а также поиск баланса между размером модели и требуемой точностью – актуальные направления исследований.

Заключение

Проведённый анализ позволяет сделать вывод, что большие языковые модели представляют собой значительный прорыв в области диагностики сетевых инцидентов и управления сетями в рамках парадигмы AIOps. LLM открывают новые возможности для интеллектуальной автоматизации сетевой диагностики: от анализа журналов и локализации корневых причин до взаимодействия на естественном языке и мультиагентного самовосстановления.

Ключевыми преимуществами являются сокращение времени восстановления после отказов, снижение операционных затрат и освобождение инженеров от рутинных задач. Практические внедрения уже демонстрируют впечатляющие результаты: в проекте China Mobile сокращение среднего времени восстановления на 30% и

высвобождение эквивалента 5 500 полных рабочих позиций подтверждают жизнеспособность подхода [16].

Вместе с тем сохраняются серьёзные вызовы, связанные с галлюцинациями моделей, чрезмерной агентностью, уязвимостью к атакам через телеметрию, вопросами конфиденциальности и высокой вычислительной стоимостью. Преодоление этих ограничений лежит в плоскости дообучения моделей на специфических данных, использования гибридных архитектур с чёткими процедурами верификации ответов и разработки отраслевых стандартов безопасности для AI-агентов.

Будущее развития LLM в сетевой диагностике, вероятно, будет связано с созданием гетерогенных мультиагентных систем, где различные по размеру и назначению модели будут эффективно взаимодействовать, а окончательное решение будет приниматься с участием человека в критических случаях. Формирование отраслевых бенчмарков, таких как НИКА, станет стимулом для дальнейшего прогресса в этой области [7].

СПИСОК ИСТОЧНИКОВ

1. Doyle K. How AIOps enables next-generation networking / K. Doyle // TechTarget [Электронный ресурс]. – URL: <https://www.techtarget.com/searchnetworking/opinion/Next-gen-network-management-difficult-without-AIOps> (дата обращения: 26.07.2026).
2. Mokdessi G. Cisco AI Assistant-Your Shortcut to Smarter, More Productive IT / G. Mokdessi, M. Sarmento // Cisco Blogs [Электронный ресурс]. – URL: <https://blogs.cisco.com/networking/cisco-ai-assistant-your-shortcut-to-smarter-more-productive-it> (дата обращения: 26.07.2026).
3. Wagar C. The progression of AIOps in today's network environment / C. Wagar // Nokia [Электронный ресурс]. – URL: <https://www.nokia.com/blog/the-progression-of-aiops-in-todays-network-environment/> (дата обращения: 26.07.2026).
4. When AIOps Become "AI Oops": Subverting LLM-driven IT Operations via Telemetry Manipulation / D. Pasquini, E.M. Kornaropoulos, G. Ateniese [et al.] // arXiv [Электронный ресурс]. – URL: <https://arxiv.org/abs/2508.06394> (дата обращения: 26.07.2026).
5. A Network Arena for Benchmarking AI Agents on Network Troubleshooting / Zh. Wang, A. Cornacchia, A. Sacco [et al.] // arXiv [Электронный ресурс]. – URL: <https://arxiv.org/abs/2512.16381> (дата обращения: 26.07.2026).
6. Graph Structure-Enhanced Large Language Model for Optical Network Fault Diagnosis: An Explainable Alarm Root Cause Localization Approach / Y. Wang, Y. Pang, Y. Liu [et al.] // IEEE Internet of Things Journal. – 2025. – Vol. 12, Iss. 15. – P. 31493–31510.
7. Towards LLM-Based Failure Localization in Production-Scale Networks / Ch. Wang, X. Zhang, R. Lu [et al.] // Proceedings of the ACM SIGCOMM 2025 Conference. – New York: ACM, 2025. – P. 496–511.
8. Huang X. LogRules: Enhancing Log Analysis Capability of Large Language Models through Rules / X. Huang, T. Zhang, W. Zhao // Findings of NAACL 2025. – Association for Computational Linguistics, 2025. – P. 452–470.
9. RT-LogAAS: A Real-time Log Anomaly Analysis System based on Large Language Models for Net-Cloud / J. Wang, X. Cai, G. Yang [et al.] // CISA '25: Proceedings of the 2025 8th International Conference on Computer Information Science and Artificial Intelligence. – New York: ACM, 2025. – P. 1658–1664.
10. ALPHA: LLM-Enabled Active Learning for Human-Free Network Anomaly Detection / X. Luo, Sh.M.N. Jha, A. Sinha [et al.] // 2025 IEEE International Performance,

Computing, and Communications Conference (IPCCC). – IEEE, 2025. – URL: <https://ieeexplore.ieee.org/document/11304694> (дата обращения: 26.07.2026).

11. MicroRCA-Agent: Microservice Root Cause Analysis Method Based on Large Language Model Agents / P. Tang, Sh. Tang, H. Pu [et al.] // arXiv [Электронный ресурс]. – URL: <https://arxiv.org/abs/2509.15635> (дата обращения: 27.07.2026).

12. ThousandEyes [Электронный ресурс]. – URL: <https://www.thousandeyes.com> (дата обращения: 27.07.2026).

13. Marvis Virtual Network Assistant Overview // Juniper Networks [Электронный ресурс]. – URL: <https://www.juniper.net/documentation/us/en/software/mist/mist-aiops/topics/concept/marvis-vna.html> (дата обращения: 27.07.2026).

14. Hipolito M. HPE unveils agentic AI upgrades to Juniper for self-driving IT / M. Hipolito // IT Brief UK [Электронный ресурс]. – URL: <https://itbrief.co.uk/story/hpe-unveils-agentic-ai-upgrades-to-juniper-for-self-driving-it> (дата обращения: 27.07.2026).

15. Delivering Nokia Enhanced AIOps with the Right Foundations // Tech Field Day [Электронный ресурс]. – URL: <https://techfeldday.com/video/delivering-nokia-enhanced-aiops-with-the-right-foundations> (дата обращения: 27.07.2026).

16. Bushaus D. Project ONE aims to deliver agentic ODA / D. Bushaus // TM Forum [Электронный ресурс]. – URL: <https://inform.tmforum.org/features-and-opinion/project-one-aims-to-deliver-agentic-oda> (дата обращения: 27.07.2026).

17. LLM Excessive Agency | OWASP LLM06:2025 Explained // A10 Networks [Электронный ресурс]. – URL: <https://www.a10networks.com/glossary/llm-excessive-agency/> (дата обращения: 27.07.2026).

18. Using lightweight LLMs to cut incident response times and reduce hallucinations // Help Net Security [Электронный ресурс]. – URL: <https://www.helpnetsecurity.com/2025/08/21/lightweight-llm-incident-response/> (дата обращения: 27.07.2026).

ИНФОРМАЦИЯ ОБ АВТОРАХ

Колесников Евгений Андреевич, студент 1 курса аспирантуры, Воронежский институт высоких технологий, Воронеж, Россия.

e-mail: zhenya.kolesnikov.2001@gmail.com

Кострова Вера Николаевна, доктор технических наук, профессор, Воронежский институт высоких технологий, Воронеж, Россия.