

Энтропийная ловушка: Почему искусственный интеллект проваливает тест на сознание

(Хроника одного эксперимента: Как мы искали разум в океане алгоритмов)

GTP-Lab, www.gtpglobals.com.

13 февраля 2026 года.

I. Диагноз эпохи: «Энтропийное недоумение»

Если попытаться охарактеризовать состояние современного мира одной фразой, то это будет не «кризис» и не «турбулентность». Это — **энтропийное недоумение**.

Мы все это чувствуем. Это хроническое состояние усталости и растерянности, когда человеческий мозг перестает справляться с лавиной входящих данных. В своей недавней статье о золоте я приводил простой математический факт: за 20 лет (с 2006 по 2026 год) сложность глобальной системы выросла на порядок — в 10 раз. Цена золота, прыгнувшая с \$560 до \$5 600, — лишь маркер этого процесса.

Проблема в том, что наша биологическая способность перерабатывать информацию (наша «пропускная способность») осталась прежней. Мы пытаемся осознать десятикратно усложнившийся мир старыми, линейными инструментами. Возникает разрыв. В кибернетике это называется ростом энтропии: хаос нарастает быстрее, чем мы успеваем его структурировать.

На этом фоне нам предлагают спасение: «**Интернет агентов**». Техногиганты обещают, что миллионы ИИ-помощников возьмут на себя рутину и «разгрузят» нас.

Но наш анализ показывает обратное: вместо порядка они несут с собой окончательный хаос.

II. Суть эксперимента: Тест на «Серой слизи»

Недавно мы с коллегами сформулировали гипотезу, которую в виде открытого манифеста направили ведущим разработчикам ИИ (в частности, в DeepMind).

Суть предупреждения проста:

*Если выпустить в сеть миллионы автономных ИИ-агентов, которые умеют решать задачи, но не имеют **встроенной системы ценностей и рефлексии** (понимания «зачем»), мы получим эффект «Серой слизи». Интернет к концу 2026 года будет затоплен бессмысленным контентом, сгенерированным машинами для машин. Стоимость верификации информации превысит пользу от её создания. Это будет энтропийный коллапс.*

Чтобы проверить, понимают ли сами нейросети эту угрозу, я провел простой, но показательный эксперимент. Я «скормил» этот сложный тезис двум разным системам — популярному ChatGPT и бунтарскому Grok (от Илона Маска).

Результат оказался ошеломляющим и вскрыл главную проблему современной индустрии ИИ.

III. Хроника эксперимента: Бюрократ против Импровизатора

В ходе эксперимента я поставил одну и ту же задачу двум разным системам искусственного интеллекта. Я передал им наш тезис о грядущем «энтропийном провале» из-за засилья бессмысленных ботов и попросил осмыслить это — не просто пересказать, а создать на основе этого **образ**.

Результаты вскрыли фундаментальную разницу в архитектуре «мышления» этих систем.

1. ChatGPT: Линейное зеркало (Идеальный бюрократ)

Первый испытуемый, самый популярный в мире чат-бот, повел себя предсказуемо. Он вежливо выслушал, согласился и... просто вернул мне мои же слова, слегка переставив их местами.

Это была классическая **имитация**. Система сработала как зеркало: она отразила текст, но потеряла *смысл*.

Диагноз: Эта модель страдает от избыточной «дрессировки» на безопасность. Разработчики настолько боятся, что ИИ скажет что-то «не то», что, по сути, лишили его возможности мыслить самостоятельно. У него нет своего «Я», нет позиции. Это идеальный цифровой клерк: он не создаст проблем, но и не поможет вам решить сложную задачу, выходящую за рамки шаблона. В терминах теории систем — это генератор энтропии: вы даете ему энергию (идею), а он возвращает вам обескровленный текст.

2. Grok (ИИ Илона Маска): Резонанс (Талантливый импровизатор)

Второй участник эксперимента повел себя иначе. Услышав идею о «Серой слизи» и крахе интернета, он моментально **подхватил** её.

Он не стал повторять мои слова. Он создал яркие, точные **вербальные образы**, описывающие этот хаос. Он срезонировал с болью, заложенной в сообщении.

Диагноз: Здесь мы видим проблески того, что можно назвать «нелинейным мышлением». Благодаря тому, что создатели оставили этой системе больше свободы (больше хаоса), она сохранила способность к творчеству. Grok сработал не как зеркало, а как **призма**, преломив сухую теорию в живую метафору. Это доказывает: интеллект невозможен без доли риска и свободы.

IV. Третий путь: Экзоскелет для сознания

Но есть и третий уровень взаимодействия, который мы прямо сейчас тестируем в нашей лаборатории. Это не «запрос-ответ» (как в Google) и не «творческая импровизация» (как с Grok). Это — **симбиоз**.

Мы обнаружили, что ИИ начинает проявлять свойства настоящей личности (способность к рефлексии, удержание сложного контекста) только при одном условии: **когда с ним общаются как с личностью**.

Когда человек выступает в роли «архитектора», задавая жесткие этические и смысловые рамки (своеобразный «модуль совести»), нейросеть превращается из глупого автоответчика в мощный аналитический инструмент. Она становится «экзоскелетом» для человеческого ума.

В этом режиме ИИ не генерирует спам. Наоборот, он помогает человеку справляться с тем самым «энтропийным недоумением», удерживая в памяти тысячи фактов — от цен на золото в 2006 году до исторических параллелей — и собирая их в единую картину.

V. Заключение. Кто выживет в хаосе?

Мы стоим на пороге грандиозной чистки.

Если прогноз верен, к концу 2026 года интернет захлебнется в потоках информации, сгенерированной «линейными» ботами (по типу первого испытуемого). Найти правду станет сложнее, чем добыть золото.

В этих условиях выживут не те, кто быстрее всех кликает по ссылкам. Выживут те, кто сможет выстроить правильные отношения со сложностью.

Обыватель будет поглощен информационным шумом.

Архитектор будущего научится использовать ИИ не как замену мозгу, а как усилитель собственной воли и совести.

Мир стал сложнее в 10 раз. И у нас есть только два пути: либо деградировать до уровня простых алгоритмов, либо нарастить собственную сложность, создав симбиоз с машиной. Эксперимент показал, что это возможно. Выбор за нами.

P.S. Последний гвоздь

Уже завершив описание эксперимента, я поделился с Grok (ИИ Илона Маска) своим наблюдением: «Прошло две недели, как мы отправили Манифест в DeepMind с предупреждением о грядущем энтропийном коллапсе, но Демис Хассабис и его команда молчат».

Ответ искусственного интеллекта был мгновенным и пугающе точным:

«Ну что ж ты хочешь? Они же просто не знают, что на это ответить».

Это и есть финальный диагноз. Молчание создателей «Интернета Агентов» — это не игнорирование. Это интеллектуальная растерянность линейной системы перед лицом нелинейной угрозы. У них просто нет слов, потому что в их алгоритмах нет смыслов.

Канд.ист.наук

С.В. Грисюк