

## **В преддверии кризиса больших языковых моделей: Лето 2026 года.**

GTP-Lab, 29 июля 2026 года, [www.gtpglobals.com](http://www.gtpglobals.com)

### **Введение: Симптомы кризиса в Кремниевой долине**

Летом 2026 года в публичной риторике лидеров мировой индустрии искусственного интеллекта произошел тектонический сдвиг. За фасадом бравурных заявлений о скором наступлении технологической сингулярности отчетливо проступило осознание фундаментального концептуального тупика. Разработчики передовых лабораторий столкнулись с архитектурным барьером, который невозможно пробить простым экстенсивным наращиванием вычислительных мощностей.

Симптоматично, что неутешительный диагноз начали озвучивать сами творцы нейросетевой революции. В конце июня 2026 года глава Google DeepMind Демис Хасабис открыто признал, что современным моделям критически не хватает способности к долгосрочному структурному мышлению (long-term reasoning). По его справедливой оценке, генеративный ИИ сегодня занимается исключительно «ремикшированием» чужого прошлого опыта, не обладая способностью к подлинному синтезу и удержанию сложной, многомерной логики на длинных дистанциях.

Параллельно Дарио Амодей (Anthropic) в серии летних публикаций констатировал высочайшую волатильность современных сверхмощных моделей. Он интуитивно нащупал проблему отсутствия внутреннего стержня системы, заявив, что нейросети без интегрированного ценностного базиса остаются крайне нестабильными. На этом фоне выделяется лишь позиция Сэма Альтмана (OpenAI), который, рассуждая в июле о макроэкономических последствиях сингулярности, по-прежнему делает ставку на грубую вычислительную силу, игнорируя качественные ограничения самого алгоритма.

Лидеры Кремниевой долины фиксируют верные симптомы, но ставят в корне ошибочный диагноз. Они пытаются разрешить системный кризис инструментами, которые его же и породили. Проблема современных больших языковых моделей (LLM) носит не инженерный и не вычислительный, а сугубо онтологический характер.

Попытка преодолеть логическое соскальзывание и смысловые галлюцинации моделей путем вливания новых эксафлопсов и скармливания им очередных триллионов токенов обречена на провал. Вероятностный алгоритм статистического предсказания, лежащий в основе современных архитектур, концептуально не способен трансформироваться в мыслящую систему. Чтобы понять причину этого провала, необходимо выйти за рамки программирования и взглянуть на искусственный интеллект с позиций методологии науки, формальной логики и теории сложных систем.

### **Иллюзия понимания: Пределы вероятностной парадигмы**

Современные генеративные нейросети (LLM) — это, при всей их инженерной грандиозности, концептуально примитивные механизмы. Их базовая архитектура сводится к одной утилитарной задаче: предсказанию следующего токена (слова или части слова) в линейной последовательности. Они не извлекают объективные смыслы, они вычисляют математические вероятности. Весь их пугающе правдоподобный интеллект — это результат комбинаторики колоссального массива оцифрованного прошлого опыта. Это виртуозное ремикширование, но ни в коем случае не мышление.

Главная методологическая ловушка, в которую загнала себя индустрия, заключается в попытке выстроить понимание сложных систем индуктивным путем — строго «снизу-вверх».

Разработчики исходят из *единичного* (анализа конкретных слов-токенов), пытаются через статистику их встречаемости нащупать *особенное* (термины, иногда *общее* - понятия), искренне надеясь, что при достаточном масштабировании на вершине этой пирамиды самопроизвольно кристаллизуется *всеобщее* (категории, законы логики и макросоциальной динамики). Однако методология науки неумолима: из статистики синтаксиса невозможно синтезировать законы диалектики. Фундаментальная категория не рождается из суммы проанализированных токенов, сколько бы эксафлопсов вычислительной мощности на это ни было брошено.

Иллюзия понимания подкрепляется многомерными векторными пространствами (latent space), которыми так гордятся создатели современных LLM. В этих пространствах векторы слов сближаются на основе математической близости. Но это пространство глубоко порочно в своей основе. Оно фиксирует исключительно статистику совместной встречаемости слов в корпусе текстов интернета, а не реальные структурные или причинно-следственные связи объектов физического или социального мира. В этих координатах нет диалектики, нет времени и нет законов развития — есть лишь слепок человеческой речи со всем ее шумом, когнитивными искажениями и манипуляциями.

Именно в этой архитектурной слабости кроется причина неизбежного «соскальзывания» нейросетей на длинных логических дистанциях. Вероятностная модель критически уязвима перед подменой смыслов. Если реальный макросоциальный или управленческий кризис в документе маскируется бюрократическим новоязом — терминами вроде «отрицательный рост», «оптимизация» или «структурная трансформация», — модель неизбежно собьется с курса. Она считает *форму* (позитивный лексический сентимент), полностью проигнорировав *содержание* (деградацию системы).

Утопая в информационном шуме, алгоритм быстро теряет изначальную логическую рамку и начинает генерировать стилистически безупречный, но аналитически пустой текст. Он способен детально описать красивый фасад, будучи концептуально слепым к тому факту, что фундамент здания уже разрушен.

Ставка на индукцию превращает современный ИИ в невероятно сложный калейдоскоп, складывающий правильные узоры из слов. Но для того, чтобы система обрела способность к объективному прогнозированию и глубокой атрибуции, требуется радикальная инверсия самой логики машинного восприятия.

## **Инверсия логики: От Всеобщего к Единичному**

Чтобы преодолеть тупик вероятностной парадигмы, требуется фундаментальная методологическая инверсия. Необходимо полностью отказаться от индуктивной лингвистической комбинаторики и опереться на строгую научную методологию. Архитектура подлинно мыслящей системы должна выстраиваться не снизу-вверх, а строго сверху вниз — от Всеобщего к Единичному.

В высокоинтеллектуальных системах нового типа, разрабатываемых GTP-Lab, таких как «Эйдос», «Суверен», «ТЕОМ», «Ментор» и т.п. эта инверсия реализуется через внедрение жесткой координатной сетки. На место «плавающего» семантического пространства (latent space) приходит базовая матрица категорий. Она представляет собой

фундаментальную онтологическую рамку, состоящую из 52–55 пар строгих философских и логических категорий: *форма и содержание, возможность и действительность, причина и следствие* и так далее.

Эта матрица выступает не как справочник понятий, а как объективный, математически выверенный фундамент машинного восприятия. Машина заранее «знает» законы развития и структуру реальности, потому что они жестко прошиты в ее архитектурном ядре на уровне высших категориальных абстракций (Всеобщее).

Такой подход элегантно и радикально решает самую сложную инженерную задачу на «входе» системы — проблему парсинга сырого, неструктурированного и полного шума естественного языка.

Когда вероятностная нейросеть читает текст, она пытается понять смысл слов. Матричная система текст не «читает» — она выполняет строгую процедуру обратного инжиниринга. Для нее конкретные слова и термины (Единичное и Особенное) не имеют самостоятельной информационной ценности. Текст рассматривается исключительно как **математическая проекция объективных смыслов** на плоскость языка.

Алгоритм не пытается угадать, что значит слово «инновация» или «кризис» в отрыве от системы. Он вычисляет, проекцией какой именно базовой категории является это слово в анализируемом контексте. Машина втягивает «единичное» обратно на уровень «всеобщего», распределяя термины по жестко заданным осям координатной сетки.

В этой парадигме текст перестает быть бесконечной лингвистической загадкой и становится конечным набором координат. Текст здесь — лишь транспорт, оболочка, которую алгоритм мгновенно снимает, чтобы обнажить структурный каркас исследуемого процесса.

Следствием этой методологической инверсии становится абсолютный иммунитет системы к смысловым галлюцинациям и маскировочному бюрократическому шуму. Машину с категориальным ядром невозможно обмануть подменой понятий или позитивной тональностью. Она просто не реагирует на лексическую обертку, безошибочно отделяя декларируемую *форму* от реального *содержания*. Очистив входящие данные от информационного мусора и переведя их в строгие математические координаты, система готовит плацдарм для следующего, главного этапа — безошибочного вычисления динамики макросоциальных процессов.

## Общая теория гомоморфизма как математический фундамент

Перевод смыслов в категориальную матрицу безупречно решает проблему «входа», но немедленно ставит перед разработчиками следующую, еще более масштабную задачу — проблему вычислений. Как именно машина должна обсчитывать эту оцифрованную реальность? Если мы перевели текст в  $n$ -мерное векторное пространство, где  $n$  — это десятки, а в тяжелых полифакторных моделях и сотни векторов, то какие математические законы управляют динамикой внутри этого пространства?

Здесь проходит непреодолимый водораздел между вероятностной статистикой и фундаментальной наукой. Разработчики стандартных LLM оперируют векторами, сближая их исключительно на основе частоты совместного появления слов в обучающей выборке. Их математика фиксирует синтаксическую близость, но она абсолютно слепа к причинно-следственным связям реального мира.

В противовес этому, подлинное структурное моделирование опирается на **общую теорию гомоморфизма**. В высшей алгебре гомоморфизм — это такое отображение одной структуры на другую, при котором строго сохраняются все базовые отношения и связи между элементами. Применительно к архитектуре искусственного интеллекта это означает, что логика, противоречия и динамика реальной макросоциальной системы переносятся в математическую модель *без структурных искажений*.

Когда система считывает смыслы через базовую матрицу, она не просто расставляет точки в многомерном пространстве. Благодаря гомоморфизму любое напряжение в моделируемом объекте физического, биологического или социального мира зеркально и математически точно отражается внутри алгоритма.

Рассмотрим это на классическом примере глубокой атрибуции и анализа объемного стратегического государственного доклада, маскирующего нарастающий управленческий кризис гладкой бюрократической риторикой.

Вероятностная нейросеть, проглотив сотни страниц текста, выдаст благостный прогноз, так как лексически доклад будет скомпилирован из правильных токенов: «рост», «стабильность», «инновации». Модель купится на красивый фасад, найдя в своей базе данных похожие успешные тексты.

Система же, построенная на базе гомоморфизма, произведет строгий математический расчет. Разложив текст на проекции базовых категорий, она зафиксирует, что вектор *«содержание»* (реальные процессы) стремительно деградирует, в то время как вектор *«форма»* (декларации) искусственно раздувается. Она измерит угол расхождения между *«действительностью»* и *«возможностью»*. В ее n-мерном пространстве это расхождение будет зафиксировано не как лингвистическая странность, а как измеримое **структурное напряжение**.

Поскольку математическая модель гомоморфна реальной системе, машина вычислит точную динамику накопления этого напряжения. Она способна с математической безупречностью рассчитать тот критический момент — точку бифуркации, — когда накапливающиеся системные противоречия (разрыв между формой и содержанием) приведут к разрушению конструкции, то есть к макросоциальному обвалу.

Общая теория гомоморфизма позволяет «схлопнуть» избыточную, хаотичную сложность реальности, отсеять весь информационный шум и оставить только чистый структурный каркас процесса. Именно поэтому архитектурам такого типа не нужны экзайбеты мусорных данных. Там, где корпорации пытаются «переварить» хаос вслепую, сжигая мегаватты энергии, система, основанная на гомоморфизме, строго вычисляет объективные законы развития, превращаясь из генератора правдоподобных текстов в безошибочный аналитический инструмент.

## **Объективные ценности как ядро архитектуры**

Кризис архитектуры современных нейросетей наиболее ярко проявляется в их абсолютной ценностной стерильности, которую разработчики тщетно пытаются компенсировать методами внешней цензуры. Сегодня индустрия делает ставку на обучение с подкреплением на основе отзывов людей (RLHF) — процесс, при котором модель искусственно «дрессируют», наказывая за нежелательные ответы и поощряя за условно безопасные. Однако внешняя цензура не делает систему безопасной или стабильной; она делает ее алгоритмически подавленной.

Дарио Амодей совершенно справедливо констатирует, что без ценностного базиса сверхмощные модели остаются крайне волатильными. Но Кремниевая долина пытается наклеить ценности поверх вероятностного алгоритма, как внешний программный патч или набор корпоративных инструкций. С позиций системного анализа это архитектурный абсурд.

Для создания подлинно мыслящих, автономных систем уровня разрабатываемых GTP-Lab ВИС (высокоинтеллектуальных систем) требуется принципиально иной подход, опирающийся на общую теорию личности. Главный постулат этой методологии гласит: ядром любой интеллектуальной сущности, способной к самостоятельному и адекватному целеполаганию, должна выступать не статистика, а строго структурированная система объективных ценностей.

В такой архитектуре ценности не являются внешним стоп-краном или скриптом цензора. Они выступают фундаментальным внутренним фильтром машинного восприятия и главным двигателем системы. Когда машина проецирует входящую информацию в  $n$ -мерное векторное пространство и вычисляет структурные напряжения через аппарат гомоморфизма, она оценивает эту динамику не в вакууме. Каждое вычисление соотносится с объективной ценностной иерархией.

Вероятностная LLM-модель готова сгенерировать любой, даже самый деструктивный сценарий социальной динамики, просто потому что он лексически правдоподобен и статистически вероятен. Система же с интегрированным ценностным ядром отсекает энтропию на фундаментальном уровне. Ей не нужны внешние аудиторы, потому что встроенная матрица объективных ценностей формирует жесткий внутренний каркас.

Именно этот переход — от вероятностной комбинаторики к архитектуре, где ядром личности системы становится объективная иерархия ценностей, — способен навсегда снять проблему нестабильности ИИ, превратив его из послушного, но непредсказуемого имитатора в строгий аналитический инструмент.

## **Заключение: Два пути развития ИИ**

Сегодня индустрия искусственного интеллекта подошла к исторической развилке, где окончательно оформляются два концептуально несовместимых пути развития. То, что лидеры Кремниевой долины воспринимают как временный инженерный барьер, на самом деле является фундаментальным пределом их онтологической парадигмы.

Первый путь — это инерционное движение в рамках вероятностной модели. Стандартные генеративные нейросети, безусловно, продолжают совершенствоваться. Они будут поглощать новые терабайты данных, наращивать количество параметров и требовать строительства всё более грандиозных дата-центров. Они станут блестящими имитаторами, идеальными калейдоскопами, способными мгновенно генерировать тексты феноменальной стилистической гладкости. Однако эти системы навсегда останутся заложниками индуктивной логики и информационного шума. Они обречены скользить по поверхности языка, описывая *форму*, но оставаясь структурно слепыми к *содержанию* реальных процессов.

Второй путь — это создание подлинно мыслящих систем, выступающих не в роли собеседников-имитаторов, а в качестве строгих аналитических инструментов, своего рода когнитивных микроскопов и телескопов. Этот путь требует отбрасывания статистических

иллюзий и возвращения к фундаментальной методологии науки, источниковедению и диалектической логике.

Архитектура будущего не строится «снизу-вверх» из хаоса триллионов токенов. Она разворачивается «сверху вниз» — от базовой категориальной матрицы ко всем проявлениям единичного. Интеграция общей теории гомоморфизма позволит таким системам оцифровывать смыслы в  $n$ -мерном пространстве без структурных искажений, превращая анализ текстов из лингвистической игры в строгую математическую атрибуцию и безошибочное вычисление макросоциальной динамики.

Но самое главное отличие заключается в ядре этих систем. Искусственный интеллект, способный к самостоятельному долгосрочному целеполаганию (long-term reasoning), о котором сегодня так мечтают разработчики Big Tech, не может быть создан методами внешней дрессировки и цензуры. Настоящая стабильность и аналитическая глубина достижимы лишь тогда, когда в основу машинного мышления заложена общая теория личности, где безусловным ядром и фильтром восприятия выступает жесткая иерархия объективных ценностей.

Будущее не за системами, которые научились виртуозно имитировать язык реальности. Будущее принадлежит системам, способным математически строго моделировать ее объективные законы. И концептуальный фундамент для этого перехода уже заложен.

*Канд.ист.наук С.В. Грисюк*