

После манифеста: как мы учимся честно думать о результате

Публичное эссе · Protopia Garden

Татьяна Гончаренко

Алексей Константинов

Protopia Garden · Июнь 2026

С чего начать этот разговор

В манифесте мы описали попытку собрать организацию нового типа — такую, которая способна работать с трудными общими проблемами не хуже, а лучше, чем это получается у старых институтов. Мы сказали тогда, что находимся в полушаге от первого живого запуска и собираем то, что должно стоять под такой организацией: инструменты, ритмы работы, способ хранить память, проверять себя и действовать.

Этот текст — рассказ о первом собранном куске нашей нынешней сборки. Мы собрали первый элемент того, что можно назвать архитектурой мышления будущей организации, — не ту её часть, которая действует в поле, а ту, которая думает. Этот инструмент помогает работать с самым трудным и самым запущенным вопросом всей социальной работы. Вопрос этот звучит просто: **откуда мы знаем, что наша работа на самом деле к чему-то приводит?**

Чтобы объяснить, что мы построили и зачем, придётся сначала спокойно разобраться, почему этот простой вопрос на самом деле такой трудный, и почему сектор, который призван помогать обществу, отвечает на него всё хуже.

Как организации обычно думают о результате

Представим обычную благотворительную или общественную организацию. Она хочет сделать что-то хорошее: помочь детям лучше учиться, поддержать людей после катастрофы, оживить угасающий городок. Чтобы получить на это деньги, она пишет заявку. В заявке иногда даже есть раздел, который на профессиональном языке называется «теория изменений».

Теория изменений — это, если совсем просто, история о том, как наши действия приведут к хорошему результату. Звучит она примерно так: «если мы проведём тренинги для учителей, учителя станут преподавать лучше, и тогда дети будут лучше учиться». Логичная цепочка: сделаем вот это — получим вот то. Под неё формируется бюджет, потом проводятся мероприятия, потом пишется отчёт, что мероприятия проведены.

В этой схеме нет ничего глупого. Долгие годы она работала как могла. Проблема в том, что красивая история о будущем результате и сам результат — это очень разные вещи. История гладкая, логичная, её приятно читать донору. А реальность сложная и неровная. Учителя могут пройти тренинг и ничего не изменить в классе. Дети могут начать учиться хуже по причинам, не имеющим к школе никакого отношения. Городок может угасать не потому, что в нём мало рабочих мест, а потому, что люди перестали верить, что у этого места есть будущее, — и тогда любые рабочие места не помогут.

Самое неприятное вот что. Написать убедительную теорию изменений теперь очень легко. Любая языковая модель сочинит вам гладкую, логичную цепочку за минуту. И поэтому гладкость текста больше ничего не говорит о том, понимает ли организация на самом деле, что она делает. Раньше хорошо написанная заявка была хоть каким-то признаком, что за ней стоит живая мысль. Сегодня это не так. Мы подробно писали об этом в манифесте: документ перестал быть надёжным носителем понимания.

И тут открывается то, что мы считаем главным и самым запущенным местом во всей этой истории.

Самое трудное место, которое все проскакивают

Когда появился доступный искусственный интеллект, сектор бросился автоматизировать ровно то, что и так было слабым: написание заявок, отчётов, презентаций. То есть ускорять производство тех

самых гладких документов, которые и так оторвались от реальности. Это понятно по-человечески — это самая заметная, самая утомительная рутина. Но это лечение симптома.

А по-настоящему трудное место осталось нетронутым. Это не написание документов. Это **честное мышление о результате до того, как ты начал действовать**. Думать о причине и следствии — о том, что на самом деле приводит к чему, — оказывается очень тяжело. Это требует признать, что ты многого не знаешь. Это требует удержать в голове сразу несколько объяснений того, что происходит, и не вцепиться в первое удобное. Это требует сверить свою догадку с тем, что человечество уже выяснило по этому поводу. И это требует мужества сказать вслух: вот здесь мы не уверены, а вот это мы пока просто не имеем права утверждать.

Большинство организаций чувствуют себя в этом месте крайне неуютно. И, признаться, у них почти нет на это ни времени, ни инструментов. Маленькая команда под операционным давлением физически не может сесть и сопоставить свой опыт с мировым корпусом исследований. Поэтому это место проскакивают. Его заполняют шаблонными фразами, и все участники по инерции делают вид, что за фразами стоит понимание.

Отсюда наша главная мысль, и она простая. Обычная организация начинает с заявки, бюджета и красивой истории про будущий результат. Мы считаем, что начинать надо раньше и с другого конца. Сначала — всерьёз и подробно разобраться в той самой вещи, которую инициатива собираешься поменять. Понять, как она устроена, кто в ней действует, какие там силы, что в ней вообще может измениться, а что нет. И только потом, на основе этого понимания, говорить о вмешательстве, о плане и о результате. Мы не ускоряем старую машину. Мы строим другой первый шаг.

Что именно мы собрали — и почему это не просто программа

Прежде чем рассказывать, как это работает, важно сказать, что это вообще такое. Мы собрали не программу, которая делает работу за людей, и не очередного «умного помощника», пишущего за вас тексты. Мы собрали архитектуру того, что в манифесте называли соединением человека и искусственного интеллекта в одно общее мышление.

Объяснить это можно так. У человека и у машины разные сильные стороны. Машина умеет держать в памяти огромный объём, быстро ориентироваться в большом корпусе исследований, удерживать сразу несколько версий происходящего и сравнивать их, замечать противоречия, не уставать. Человек умеет совсем другое: понимать живой контекст, чувствовать, где слова расходятся с делом, нести ответственность, принимать решение там, где готового правильного ответа нет, видеть в собеседнике человека, а не строчку данных. Если соединить их по уму, получается мышление сильнее, чем у каждого по отдельности. Не человек вместо машины и не машина вместо человека, а единый рабочий контур, у которого просто больше памяти, больше внимания и больше честности.

Поэтому когда дальше в этом тексте мы говорим «машина» или «система», мы имеем в виду именно это соединение, а не самостоятельную программу, действующую помимо людей. Машинная часть здесь — не хозяин, а служебный участник: она готовит, ищет, сопоставляет, подсказывает и, что важно, не даёт человеку незаметно срезать угол. Но выбирает, решает и отвечает всегда человек. Машина предлагает — человек решает. Эту границу мы держим жёстко: есть работа, где помощь машины уместна, и есть места, где исчезновение человека разрушает саму суть дела — встреча с чужой болью, обещание, политическое решение, ответственность за ошибку.

И вот внутри этого соединения машинная часть устроена особым образом. Не для того, чтобы упростить человеку жизнь, а наоборот — чтобы не дать ему обмануть самого себя. Дальше — о том, как именно.

Что значит «разобраться в том, что ты меняешь»

Здесь нужно объяснить одно понятие, потому что оно у нас центральное. Мы называем то, что организация собирается менять, «объектом воздействия». Звучит сухо, но что поделать, это академический подход, методология построения научного знания.

Обычная практика НГО: смотрим на как-то сформулированную проблему через призму возможных активностей, перебираем их и пытаемся согласовать логику действий.

Мы настаиваем, что сначала необходимо смотреть на саму реальность: что это за сообщество, что это за ситуация, что это за беда — во всей её сложности, с её историей, конфликтами, с её собственной жизнью, которая шла и до нас и будет идти после. Сообщество — не пустой объект, который мы «обрабатываем». У него есть своё понимание себя, и оно имеет право спорить с нашим пониманием.

Возьмём простой и узнаваемый пример. Представим небольшой городок, из которого годами уезжает молодёжь: закрываются предприятия, пустеют школы, а люди понемногу теряют веру, что у этого места вообще есть будущее. Это та самая трудная, копившаяся годами проблема, ради которой и существует социальная работа. (Тот же пример мы подробно разбираем в манифесте — но здесь он нужен просто как наглядная иллюстрация, и читать манифест для этого не обязательно.)

Старый проект приходит в такой городок с уже готовой рамкой: «будем развивать занятость». Собирает под эту рамку данные, проводит мероприятия, отчитывается. Возможно, какая-то польза будет. Но узел проблемы он может вообще не задеть, потому что смотрел не туда: настоящая беда могла быть не в нехватке рабочих мест, а, например, в том, что люди давно перестали верить любым обещаниям.

Вот здесь наше соединение человека и машины работает иначе, чем привычный проект. Оно не даёт сразу выбрать одну рамку. Вместо этого машинная часть помогает человеку выложить рядом сразу несколько разных взглядов на то, что здесь вообще происходит. Может, дело в занятости. А может — в накопленном недоверии: люди перестали верить обещаниям, цифрам, чужакам, которые приезжают, измеряют, обещают и уезжают. А может — в том, что молодёжь не видит здесь будущего, и это уже само себя разгоняет. Это разные объяснения, и они ведут к совершенно разным действиям. Система держит их все живыми и не даёт схлопнуть к одному удобному — до тех пор, пока человек сознательно не выберет рабочую версию и не запишет, почему именно её. Если вернуться к сухому языку методологии науки, мы строим объект и предметные рамки.

Дальше по каждому объяснению машинная часть задаёт человеку неудобный вопрос: а на чём оно вообще основано? Это наша догадка, или это подтверждено серьёзными исследованиями, или это где-то сработало, но совсем в другом контексте, и неизвестно, сработает ли здесь? Чтобы помочь ответить, она помогает обратиться к мировому корпусу научных работ и готовит для команды, у которой нет своего учёного, черновую разметку: вот эта идея, похоже, сильная и проверенная; вот эта правдоподобная, но спорная; вот эта модная, но слабая; а вот эта работала в Кении, но к нашему городку, скорее всего, не применима. Последнее слово в этой разметке остаётся за человеком.

И, может быть, самое важное: система требует завести **карту незнания**. Это честный список того, чего мы про эту ситуацию не знаем, где наша модель слаба или пуста. В обычной заявке зоны незнания отсутствуют, о них просто не знаю, потому что само незнание выглядят как слабость. У нас наоборот: модель, у которой нет такой карты, просто не проходит дальше. Признание «вот здесь мы не понимаем» — это не позор, это условие, чтобы вообще возможно было двигаться дальше в слое мышления о проблеме и вариантах воздействия.

Где проходит граница, которую все размывают

Есть ещё одно различие, ради которого, по сути, и собрана ещё одна часть нашей системы. Сектор постоянно сливает в одно две совершенно разные вещи: «у нас хорошая модель» и «результат реально произошёл». Это разрыв между теорией и практикой: сама по себе теория может быть верной или ошибочной, но ее связи с практикой, с реальностью — отдельный предмет заботы.

Это не одно и то же. Можно построить прекрасную, обоснованную модель того, как тренинги для учителей улучшат учёбу детей, — и при этом в реальности дети так и не станут учиться лучше. Хорошая модель — это про качество мышления. Случившийся результат — это про то, что на самом деле изменилось в мире. Между ними лежит вся непредсказуемость реальной жизни.

Поэтому наша машина держит эти два вопроса строго отдельно и никогда не позволяет одному притвориться другим. Есть отдельное доказательство того, что модель хороша. И есть отдельное — гораздо более трудное — доказательство того, что изменение действительно случилось в поле. И пока нет второго, организации не разрешается громко заявлять об успехе. Слабое, непроверенное утверждение можно оставить для внутренней учёбы или для откровенного промежуточного разговора с донором — но его нельзя превращать в публичную победу и тем более в основание для денег. Это правило у нас не пожелание, а жёстко проверяемый запрет.

Зачем эту дисциплину доверять машине

Можно спросить: если всё держится на честности мышления и на человеческом суждении, при чём тут вообще компьютер?

При том, что её почти невозможно удержать на одной силе воли — особенно когда команда устала, а донору надо отчитаться к пятнице. Слишком велик соблазн срезать угол: пропустить карту незнания, выдать догадку за доказанное, назвать случайное совпадение результатом своей работы. Человек по доброй воле обещает себе так не делать — и под давлением делает. Поэтому мы превратили эту дисциплину в набор правил, которые проверяет машина.

Можно представить это как встроенного «контролёра». Он не проверяет, права ли организация по существу — это не его дело и не его компетенция. Он проверяет другое: не срезаны ли углы. Заполнена ли карта незнания. Держатся ли ещё живыми альтернативные объяснения. Не выдаётся ли совпадение за эффект. Не заявлено ли больше, чем доказано. Если углы срезаны, он просто не пропускает работу дальше — не как начальник, а как правило, одинаковое для всех. Так честность перестаёт зависеть только от доброй воли уставшего человека и становится встроенным свойством самого процесса.

Но «контролёр» — это лишь одна сторона. Если сказать о главном, машина здесь делает для нас две вещи. Первая — помогает дисциплинировать мышление и удержать его сложность. Человеческий ум не может одновременно держать в голове десяток объяснений, сотни исследований и всю историю прошлых решений — он неизбежно упрощает, сводит к одной удобной картинке. Машина берёт эту тяжесть на себя: помнит, сравнивает, не теряет нити — и за счёт этого мы можем позволить себе думать о проблеме сложнее, чем в одиночку, не сваливаясь в плоскую версию только потому, что её легче удержать. Вторая — помогает увидеть и обойти когнитивные заблуждения, которым подвержены все мы без исключения: склонность вцепиться в первую же версию, принять желаемое за доказанное, замечать только то, что подтверждает нашу правоту, и не замечать остального. Машина не умнее нас в этом — но она бесстрашна и не устаёт, и потому способна вовремя показать: смотри, вот здесь ты, кажется, поддаёшься привычной ошибке. Не она решает за человека — она возвращает человеку возможность увидеть себя со стороны.

Как это происходит: разговор человека и машины

Чтобы всё это не звучало пугающе, опишем, как это выглядит на практике. Никакого пульта с тумблерами и сложных анкет здесь нет.

Человек приходит с обычными словами — так, как нормальный человек описывает наболевшую ситуацию: «у нас в районе спивается молодёжь, и мы не знаем, что с этим делать»; «после наводнения люди не возвращаются в посёлок». Не на профессиональном языке, не теорией изменений — просто живым рассказом.

Дальше начинается разговор. Можно говорить голосом или писать текстом, можно приложить документы — отчёты, заметки, расшифровки интервью, всё, что есть под рукой. Машина слушает и не торопится выдать готовый ответ. Она задаёт уточняющие вопросы — те самые, которые в спешке обычно никто не задаёт: а кого именно это затрагивает? а что вы уже пробовали? а что говорят сами люди? Из этого разговора она постепенно собирает структурированные записи — мы называем их карточками: отдельный взгляд на проблему, отдельная причинная связь, отдельное утверждение, с пометкой, где догадка, а где есть основание. Когда ей нужно что-то проверить, она сама идёт в научную литературу и в интернет, приносит найденное, перепроверяет себя и прямо показывает, где подтверждения не нашла.

Получается не допрос и не экзамен, а скорее работа с очень начитанным, неутомимым и дотошным помощником: он прочитал почти всё, ничего не забывает и мягко не даёт вам проскочить трудное место. Но всё, что он собрал, остаётся черновиком до тех пор, пока человек это не посмотрел, не поспорил и не утвердил. Так шаг за шагом — от живого рассказа к проверенной модели — и рождается тот объект, о котором шла речь.

Из чего это собрано: три связанных узла

Если заглянуть внутрь, всё это держится на трёх тесно связанных частях. Назовём их прямо — чтобы было понятно, что речь про конкретный, собранный инструмент, а не про туманный «искусственный разум».

Первый узел — общая база знаний. Раз человек и машина работают в связке, им нужна общая память, к которой одновременно есть доступ и у людей, и у машинной части. Это не папка с файлами и не архив отчётов, а живое организованное хранилище: сюда складываются прошлые модели, собранные доказательства, сигналы из поля, принятые решения и причины, по которым их потом пересматривали. Главное свойство такой памяти — она не уходит вместе с уволившимся сотрудником и не начинается каждый раз с чистого листа. Организация перестаёт изобретать себя заново при каждой смене команды, и её опыт становится не грудой папок, а работающим ресурсом, в который и человек, и машина могут заглянуть в любой момент.

Второй узел — модель объекта воздействия — живёт внутри этой базы знаний. Это сердцевина всей конструкции — та самая модель объекта, о которой шла речь: **машиночитаемый объект** из структурированных записей-карточек — отдельная на каждый взгляд (рамку) на проблему, на каждый причинный механизм, на каждое доказательство, на каждую гипотезу. Такие записи читает и человек, и программа. Модель версионируется, как программный код: видна вся её история — что менялось, кто добавил, почему. Из неё можно сделать разные «виды» под разные задачи — причинную схему, карту действующих лиц, привычную теорию изменений для языка доноров, — но ни один такой вид не выдаётся за единственную истину. К модели подключён **академический слой**: связь с мировыми открытыми каталогами научных работ (например, OpenAlex и Semantic Scholar). Когда надо проверить, на чём держится объяснение, машина помогает обратиться к этой литературе

— ищет работы, по возможности проверяет, не была ли статья отозвана, находит обобщающие обзоры, готовит черновик доказательной карточки и прямо помечает, чего не нашла. А суждение — годится ли это исследование для нашего случая, насколько крепко оно сделано — остаётся за человеком. И поверх всего работает тот самый «контролёр»: маленькая отдельная программа (на нашем языке — линтер), которая пробегает по записям и проверяет, соблюдены ли правила; если углы срезаны, она не пропускает работу дальше.

Третий узел — контур мандата на импакт. Если первая часть отвечает на вопрос «достаточно ли хорошо мы понимаем объект, чтобы вообще иметь право действовать?», то этот узел отвечает на следующий: «что мы имеем право утверждать о результате — и кому что обещано?». Проще говоря, мандат на импакт — это открытое и заранее проговорённое соглашение о том, что именно должно измениться в реальности, на каком горизонте, при каких ограничениях, кто вправе ставить такую цель и как мы поймём, что изменение действительно произошло.

У этого соглашения есть несколько важных свойств, и они не случайны. Оно может родиться с двух сторон: от донора, который задаёт цель и не выбирает заранее исполнителя (а дальше идёт открытый отбор тех, кто возьмётся), — или со стороны самого поля, когда инициатива приходит с реальной проблемой и моделью. Мы сознательно разрешаем второй путь: если позволять задавать цели только со стороны денег, система с самого начала встроит в себя несправедливость. Внутри мандата прямо записаны защиты от такой несправедливости: какова роль сообщества (соавтор, с ним советуются, его просто информируют), есть ли его согласие, есть ли у него право оспорить мандат ещё до того, как он запущен, и кто контролирует данные и права на результат. В чувствительных ситуациях открытое согласие может быть небезопасным — тогда мандат идёт в «узком» режиме, без громких публичных заявлений.

И главное — здесь живёт тот самый разрыв между «модель хороша» и «результат случился», о котором шла речь выше. Из него вырастает жёсткое правило: ни одно утверждение о результате не может стать публичным, финансово значимым или превращённым в какой-либо актив без обоих доказательств — и модельного, и полевого — и без отдельного решения о проверке. Слабое или одностороннее доказательство годится максимум для внутренней учёбы и спокойного промежуточного разговора с донором, но не для громкой победы. Исключений нет.

Вместе эти три узла и образуют ту самую архитектуру мышления: общая память, в которой живёт дисциплинированная модель реальности, и грамматика, которая не даёт перепрыгнуть от красивой модели к значимому утверждению о результате без доказательства и без права его сделать.

Вернёмся к вопросу, с которого мы начали: откуда мы знаем, что наша работа к чему-то приводит? Честный ответ начинается не с громкого «у нас получилось», а с этой дисциплины — понять объект, удержать сомнение, развести модель и результат и не заявить больше, чем доказано.

Модель собрана — что с ней делать дальше

Сама по себе модель — не цель и не разовая поделка. Это рабочий инструмент, и пользуются им по-разному в зависимости от того, на каком этапе вы находитесь — от самых первых смутных вопросов до уже идущего проекта. И, что важно, он не застывает: он живёт и меняется вместе с работой.

Самый ранний и самый важный момент — стратегический, когда проекта ещё нет вовсе. Команда или инициатива только собралась и хочет взяться за какую-то проблемную область или за целый клубок проблем конкретного сообщества. По сути, ей предстоит решить стратегические вопросы: какие цели выбрать, на чём сфокусироваться, какие шаги делать в каком порядке и куда направлять всегда ограниченные и недостаточные ресурсы. Именно здесь, на стратегическом уровне, работа

с моделью даёт критически важный вклад — потому что от этих ранних выборов зависит всё остальное. Архитектура работает как пространство для совместного исследования: команда в разговоре с машиной начинает разбираться, как эти проблемы устроены на самом деле: выделяет важные контексты и структурные узлы, растаскивает запутанный клубок на составляющие, ищет точки, где на ситуацию вообще можно повлиять, — причём не на одном уровне, а на разных, от отдельного человека до института и среды вокруг. Слой за слоем, по мере того как растёт собственное понимание команды, складывается целый набор возможных вмешательств — кандидатов в будущие проекты.

А дальше из этого набора можно собрать не россыпь разрозненных активностей, а связную стратегию: где у каждого действия есть внятное обоснование, видно, как оно сцеплено с другими планами и точками воздействия, продуманы возможные побочные эффекты, а заранее выстроенное наблюдение за результатами служит ещё и собственному обучению команды. И всё это — с трезвым пониманием, что ресурсов всегда меньше, чем хотелось бы, и направлять их нужно осознанно, туда, где они дадут больше всего.

Та же ранняя работа нужна и с другой стороны — со стороны крупного, системного донора, и особенно до того, как он развернёт свои программы. Прежде чем распределять средства, такой донор может той же дисциплиной разобрать проблемное поле, увидеть, где на разных уровнях лежат точки воздействия, и собрать не набор разрозненных грантов, а связную программу — с продуманными связями между её частями, ожидаемыми эффектами и встроенным наблюдением за результатами. Тогда деньги с самого начала направлены осознанно, а не подгоняются под уже объявленные приоритеты.

Когда из такого исследования вырастает конкретный проект, модель даёт то, чего обычно не хватает: внятные, обоснованные гипотезы и правдивую причинную картину того, на что вы воздействуете. Из неё легко и быстро собирается проектная заявка — но уже не выдуманная гладкая теория изменений, а заявка, за которой стоит проверенная модель, доказательства и открытая карта незнания. Проект получается честным: донору вы показываете не догадку, а дисциплинированную картину — включая то, чего пока не знаете, и то, как собираетесь это узнать.

А ещё — раз объект и предполагаемые вмешательства прорисованы без прикрас, из модели понятнее вытекает и объём нужных ресурсов: сколько и чего реально потребует работа, видно не из подгонки под бюджетную строку, а из самой картины проблемы. Это не точная смета, но куда более трезвая опора, чем привычная прикидка на глаз. И, что неожиданно, такая заявка ещё и убедительнее обычной: видно, что за словами стоит работа мысли, а не просто складный текст.

Если вы уже в процессе реализации, модель работает иначе. Она позволяет сверять ваши версии с тем, что реально происходит в поле: что подтверждается, а что нет; где нужно расширить контекст; какие новые гипотезы подсказывает сама жизнь. Показатели из модели становятся тем, что вы отслеживаете на практике, а расхождение между моделью и данными — это не неудобство, которое хочется спрятать в отчёте, а сигнал к пересмотру. Так действие постоянно возвращается обратно в мышление: поле учит модель, а уточнённая модель ведёт следующее действие.

И здесь главное: это не разовая сборка в начале, после которой модель кладётся на полку. Это живой объект, который меняется, дополняется и корректируется обратной связью на протяжении всей жизни проекта. У него есть версии и история — видно, как менялось понимание и почему. Когда реальность спорит с моделью, модель обновляется; когда команда чему-то научилась, это остаётся в общей памяти, а не только в голове у человека, который завтра может уйти. Та же дисциплина, что собрала модель, держит её честной по мере того, как она меняется. И открыто пересмотреть собственную картину здесь — признак силы, а не провала.

От каких ошибок это бережёт

Проще всего объяснить ценность через конкретные ловушки, в которые попадают даже умные и добросовестные люди. Будем держать перед глазами наш пример — городок, из которого уезжает молодёжь, — и посмотрим, как одни и те же ошибки повторяются снова и снова. Их удобно разложить на несколько семей: искажения самого мышления, ошибки причинных моделей и проверки гипотез, трудность измерить реальный эффект, ошибки документов и языковых моделей и ошибки отношения к людям. Наша система устроена так, чтобы удерживать команду от каждой.

Начнём с **искажений самого мышления** — с устройства нашего собственного ума. У каждого из нас есть встроенные умственные сокращения: в обычной жизни они выручают, но в трудной работе тихо уводят в сторону. Под давлением времени они срабатывают незаметно, и беда в том, что изнутри ошибка ощущается как ясное понимание.

Якорение на первой версии. Самая частая ловушка ума — вцепиться в первое пришедшее объяснение и больше его не отпускать. Команда решает: «городок умирает, потому что нет работы», — и дальше всё крутится вокруг рабочих мест, а настоящий узел (например, что люди давно перестали верить любым приезжим и обещаниям) так и остаётся незамеченным. Система этому мешает: она держит несколько объяснений живыми рядом и требует сознательно выбрать рабочее, записав, почему именно его и от чего мы отказываемся.

Подтверждающий уклон. Мы склонны замечать то, что подтверждает нашу правоту, и пропускать то, что ей мешает; помнить удачи и забывать провалы; а хорошее, что случилось рядом, охотно записывать на свой счёт. Особенно сильно это работает, когда мы уже вложились в идею. Уверовав, что городку нужны рабочие места, мы начинаем отмечать каждую новую открывшуюся мастерскую как подтверждение своей правоты — и не замечать, что молодые семьи всё равно собирают вещи. Поэтому система с самого начала требует завести карту незнания — честный список того, чего мы не понимаем, — и держать рядом неудобные объяснения. Гладкая картинка, в которой всё сходится, чаще всего и есть первый признак, что мы поддались искажению.

Следующая семья — **ошибки причинных моделей и проверки гипотез**. Думать о причине и следствии — отдельное и по-настоящему трудное умение, и здесь ошибаются даже опытные люди.

Корреляция, выданная за причину. После нашей программы в городок вернулись несколько молодых семей — значит, это мы? Не обязательно: может, по соседству вновь открылся завод, и вернулись бы они в любом случае. Рядом всегда действует множество причин, а простое совпадение во времени само по себе ничего не доказывает. Система требует прямо помечать, на какой ступени мы находимся: просто заметили связь — или действительно показали, что одно вызвало другое.

Гипотеза, которую нельзя проверить. «Мы возрождаем доверие в городке» звучит хорошо — но если заранее не сказать, по каким именно наблюдениям мы поймём, что доверие НЕ растёт, такое утверждение нельзя ни подтвердить, ни опровергнуть. А неопровержимая гипотеза — это уже не рабочая модель, а вера. Поэтому для каждой причинной связи система просит заранее назвать, какое наблюдение её бы опровергло, и показать, на каких доказательствах она держится.

Наивно склеенные причинные цепочки. В нашем городке это выглядит так: «наши встречи сближают соседей» → «близость возвращает доверие» → «доверие удерживает молодёжь» → значит, наши встречи удержат молодёжь. Каждое звено по отдельности звучит разумно, но взято из разного источника разной надёжности, и вывод из их сцепки рассыпается. Система не сцепляет такие цепочки молча: она показывает, откуда взято каждое звено и насколько ему можно верить, и помнит, что вывод не крепче самого слабого звена в нём.

Дальше — **трудность измерить реальный эффект**. Допустим, модель хороша и связь действительно причинная. Остаётся самое трудное — и самое часто подделываемое: доказать, что изменение реально произошло и что это наша заслуга.

«**А что было бы без нас?**» Допустим, отток молодёжи из городка замедлился. Но, может, он замедлился бы и так — из-за общего поворота в регионе? Главный вопрос измерения импакта — это вопрос о невидимом: что случилось бы, если бы нас тут не было? Без ответа на него любое «у нас получилось» повисает в воздухе. А ответить на него добросовестно почти всегда трудно: сравнить не с чем, контрольную группу не всегда можно собрать, а жизнь не ставит чистых экспериментов. Поэтому система требует различать «мы внесли вклад наряду с другими» и «мы вызвали результат» — и не выдавать первое за второе.

Обманчивая точность. Из этой трудности рождается соблазн спрятать её за одним красивым числом — свести возрождение целого городка к строке «мы создали социального эффекта на столько-то миллионов». Число выглядит солидно, но за ним обычно прячется огромная неопределённость, а единственная цифра создаёт ложное чувство, будто мы знаем точно. Правдивее — разговор о том, какие пути развития событий возможны и чего каждый стоит. Система не даёт свести сложное к обманчиво точной цифре.

Активность вместо результата. Проще всего отчитаться за то, что мы сделали: «провели двенадцать встреч с жителями городка, собрали триста человек». Но проведённые мероприятия — это ещё не изменившаяся реальность. Чтобы говорить о результате, нужно показать, что изменилось что-то наблюдаемое и проверяемое: например, молодые семьи действительно стали оставаться, а не просто отсидели наши встречи. И даже тогда остаётся вопрос, наша ли это заслуга. Система заставляет прямо называть, что именно мы утверждаем: просто провели активность, заметили перемену рядом — или действительно показали, что наблюдаемое изменение вызвано нашим вмешательством.

Ещё одна семья — **ошибки документов и языковых моделей**. Это та самая болезнь схлопнувшегося документного слоя, с которой начинался разговор, — только теперь усиленная искусственным интеллектом.

Гладкость вместо понимания. Привычные ходы отчётов и заявок устроены так, что красиво написанный текст принимают за признак качественной работы. И теперь, с языковыми моделями, за минуту можно получить безупречно складную теорию изменений про наш городок — убедительную и при этом ни разу не сверенную с тем, что там происходит на самом деле. Поэтому наша система привязывает каждое утверждение не к красоте формулировки, а к проверяемой записи: видно, на чём оно держится и держится ли вообще.

Выдуманная учёность. У языковых моделей есть опасное свойство: они умеют уверенно сочинять — выдумать несуществующую ссылку, приписать реальному исследованию вывод, которого в нём нет, выдать правдоподобную фразу за установленный факт. Если на это опереться, получится наукообразная пустота — например, уверенная ссылка на «исследование, доказавшее, что программы занятости возрождают малые города», которого либо не существует, либо оно про совсем другие города, либо давно отозвано. Наш академический слой работает иначе: он обращается к настоящим научным базам, где может — проверяет, не была ли статья отозвана, и добросовестно записывает, что нашёл и чего не нашёл, — так что авторитет труднее подделать, а пробел — замолчать.

Мода и устаревшее знание. Сегодня в моде идея, что малые города оживляют коворкинги и «креативный класс», — но мода не доказывает правоту; а данные о причинах оттока, верные для нашего городка десять лет назад, с тех пор могли устареть. То, что идею часто цитируют, помогает её найти, но не делает её истинной. Система различает «идею мало упоминают, потому что она слабая»

и «потому что она новая», помечает свежесть доказательств и требует их переоценивать.

И наконец — **ошибки власти и отношения к людям**. Они самые опасные, потому что часто выглядят как добродетель.

Сообщество как источник данных, а не как соавтор. Люди, которых касается проблема, знают о своей жизни то, чего не видит ни один приезжий эксперт. В нашем городке настоящую причину упадка — скажем, давнюю историю обманутых обещаний, которую помнят все местные, — по-настоящему знают только сами жители; без них модель её просто не увидит. Их голос — это не «данные», которые мы собираем, а источник понимания, без которого модель почти наверняка будет неполной. И именно этот голос обычно слышат меньше всего.

Чаще всего это происходит одним из двух способов. Либо сообщество превращают в источник информации: чужаки приезжают, собирают данные, уезжают и возвращаются с готовым приговором о чужой жизни. Либо «совет с сообществом» оказывается ритуалом: людей выслушали — и всё равно решили сами, по-своему. В обоих случаях люди остаются объектом, о котором думают, а не участником, который думает вместе с нами.

Есть здесь и тонкий момент власти. Даже когда мы искренне предлагаем право возразить, оно реально лишь тогда, когда доступно и тем, у кого нет нужного языка, статуса или смелости говорить вслух. Иначе «возможность оспорить» достаётся только самым уверенным, а самые уязвимые голоса — те, ради которых часто всё и затевается, — снова не слышны.

Поэтому в нашей системе люди, которых касается дело, имеют право быть соавторами модели о самих себе: участвовать в том, как она строится, и оспаривать её ещё до того, как что-то запущено, а не только постфактум. Модель строится не про людей, а вместе с людьми.

Торопливая победа. И, наконец, главное — соблазн уже через год объявить «мы возродили городок!» и собирать под это новые деньги и громкую репутацию, пока ничего ещё толком не доказано. Именно поэтому система держит раздельно «модель хороша» и «результат случился» и не позволяет превратить надежду в публичную победу или в актив, пока изменение не подтверждено независимо.

Ни одна из этих ловушек не экзотична. Это обычная физика того, как добросовестные люди под давлением времени начинают всё лучше описывать себя и всё хуже видеть мир. Машина не делает нас умнее или честнее — она просто не даёт тихо проскользнуть мимо этих ловушек и каждый раз возвращает к ним внимание.

Что уже собрано — и чего мы пока не сделали

Чтобы не смешивать одно с другим, разделим, что у нас уже есть, а чего ещё нет.

Уже собрано и работает: общая память — база знаний; рабочий контур, который строит и проверяет модель объекта; контур, который держит грамматику честных утверждений о результате; автоматические правила-проверки, которые ловят срезанные углы; и пробные прогоны на учебных примерах. Это ещё не доказательство, что мы изменили что-то в поле, — но это уже не замысел на бумаге, а работающая, исполнимая дисциплина мышления.

А теперь о границах. У нас есть дисциплина построения хорошей модели и грамматика выверенных утверждений, но пока нет работающего способа доказывать, что изменение реально произошло в поле. Это отдельная, очень трудная задача, и это несущая стена, которую нам ещё предстоит построить.

Пока её нет, наша система — это дисциплина мышления и заготовка под будущее, а не подтверждение реальных результатов.

Куда это ведёт: к чему мы на самом деле идём

Эта архитектура мышления — не цель сама по себе. Она имеет смысл только потому, что мир вокруг изменился, а старая форма помощи за ним не поспевает. Мир стал быстрее, конфликтнее, многослойнее и длиннее в своих последствиях. А привычный способ работать с общими бедами остался медленным там, где нужна скорость; непрозрачным и вертикальным там, где нужно доверие; запутанным в согласованиях там, где нужно действовать; и он раз за разом приходит поздно — часто уже после того, как предсказуемая беда стала катастрофой. И при всём этом он чаще всего не может внятно сказать, изменил ли он вообще хоть что-нибудь.

Об этих разрывах мы подробно говорили в манифесте. Если собрать их вместе, мир предъявляет несколько требований, которым старая форма отвечает плохо: потребность быстро реагировать (некоторые беды приходят не за годы, а за дни); потребность в прозрачности и горизонтальности (чтобы помощь не обязана была идти через тяжёлые донорские структуры и долгие согласования государств и комиссий); и потребность удерживать сложность (думать о проблеме на том уровне, которого она заслуживает, а не упрощать её под формат отчёта). Всё, что мы делаем, — это попытка там, где можно, соединить усилия человека и машины так, чтобы выбросить мёртвое и лишнее, не воспроизводить старую неэффективность, удерживать настоящую сложность задачи и при этом честно отчитываться о результате. Описанная здесь архитектура мышления — тот кусок именно этого, что мы уже собрали.

Куда это выльется дальше, мы и правда пока не знаем. Одна из наших гипотез: возможно, со временем это вырастет в то или иное платформенное решение — место, где те, кто хочет изменений, и те, кто готов их поддержать, находят друг друга напрямую, вместе строят проверяемую картину происходящего и как-то собирают на это ресурсы — быстрее и прозрачнее, чем сегодня. Может быть, через неё можно будет и направлять средства, может быть, нет; может быть, для этого пригодится новая цифровая инфраструктура, а может, и не понадобится. Мы не делаем вид, что знаем форму этого решения. Это далёкий горизонт и пока только контур.

Чтобы направление стало нагляднее, представим ситуацию. Где-то в Португалии — лесной пожар; или наводнение, или внезапный наплыв беженцев. Общество готово откликнуться, но обычно все ждут государства, а оно разворачивается медленно, и между готовностью помочь и реальной помощью открывается провал, в котором гибнет время и силы. А теперь вообразим другой ход событий: те, кто готов поддержать, откликаются сразу; проблема быстро прогоняется через проверенные модели — не ради красивого отчёта, а чтобы разобраться, что на самом деле происходит, — и раскладывается на проверяемые задачи; практик на земле берётся за задачу, делает её, даёт подтвердить результат. И вся цепочка — от готовности помочь до подтверждённого реального изменения — не рвётся в том самом провале между диагнозом и действием. Мы не утверждаем, что точно знаем, как такое построить. Но именно в эту сторону нам хотелось бы двигаться.

И вот здесь видно, почему мы начали с мышления, а не с денежных рельсов. Любой такой быстрый и горизонтальный механизм — это ровно то место, куда опаснее всего пускать ложное или плохо обоснованное утверждение: скорость и автоматика только сделают ошибку крупнее и долговечнее. Поэтому первым делом нужно строить не рельсы для ресурсов, а слой, который решает, какие утверждения о результате вообще достойны того, чтобы по ним что-то двигалось. Этот слой мы и собираем.

А цифровые инструменты, которые могли бы всё это нести, — включая смарт контракты, — для нас далёкий и пока эскизный горизонт: они имеют смысл только после состоятельной модели, доказательства и проверки и поверх них, а не вместо них. Иначе говоря, прежде чем результат станет строкой в реестре, обязательством под смарт-контрактом или поводом для выплаты, он должен пройти весь путь: модель объекта, доказательную проверку, подтверждение, что изменение действительно случилось, и отдельное решение о верификации. Этот недостающий пазл большой инфраструктуры — то, что мы и собираем.

Будем честны до конца: этой платформы у нас нет, самый трудный её кусок — доказательство, что изменение действительно произошло, — ещё впереди, и мы не уверены даже, что всё это получится. Но направление нам понятно, и мы заложили необходимый краеугольный камень — без которого всё остальное превратилось бы в красивую, но опасную надстройку над пустотой.

Зачем мы это рассказываем

Мы собрали архитектуру мышления — и теперь нам нужно самое важное: заинтересованные потенциальные коллеги и партнеры. И мы будем рады собеседникам с двух разных сторон.

Со стороны поля и реализации — тем, кто прямо сейчас работает как некоммерческая организация или инициатива социальных изменений, видит, что поле вокруг изменилось, и хочет перестроить под него свою практику. Если вам близка эта архитектура мышления и вы хотите пересобрать с её помощью собственный способ думать о проблеме и о результате — мы с удовольствием сделаем это вместе с вами.

Со стороны ресурсов и поддержки — донорским структурам, фондам, программам корпоративной социальной ответственности — всем, кто хочет вкладывать средства в прозрачные и проверяемые схемы работы, а не в красивые отчёты так, чтобы средства шли под проверяемые модели и подтверждаемый результат.

И в том, и в другом случае нам интересны не те, кто ищет «искусственный интеллект, чтобы быстрее писать заявки», а те, кто чувствует ту самую боль, с которой начался этот текст: ощущение, что мы вкладываем силы, пишем гладкие отчёты — но честно не знаем, меняем ли мы что-то на самом деле. Если это чувство вам знакомо, то мы будем рады с вами познакомиться и поделиться нашими наработками.

Как и манифест, этот текст — одновременно сигнал и приглашение. Сам наш инструмент мы готовы показывать, разбирать и обсуждать — и будем рады, если его попробуют оспорить и сломать на прочность. Потому что машина, которая всё подтверждает и ничему не сопротивляется, не стоит ничего.