

УДК 159.9

DOI: 10.34670/AR.2026.29.68.024

Иллюзия объяснительной глубины на материале собственного замысла: самодиагностика в группе с помощью веб-игры

Агамалиев Рустам

Магистр,

Национальный исследовательский университет «МИЭТ»,
124498, Российская Федерация, Москва, Зеленоград, площадь Шокина, 1;
e-mail: rustam.agamaliiev@gmail.com

Дорофеев Максим

Магистр,

Национальный исследовательский университет «МИЭТ»,
124498, Российская Федерация, Москва, Зеленоград, площадь Шокина, 1;
e-mail: maxim.dorofeev@gmail.com

Аннотация

В статье проверяется воспроизводимость эффекта иллюзии объяснительной глубины (иллюзии понимания, когда человеку кажется, что он понимает предмет глубже, чем на самом деле) на материале собственного замысла участника, а не внешнего технического устройства, как в классическом исследовании Розенблита и Кейла, где испытуемые переоценивали понимание работы бытовых приборов. Экспериментальное вмешательство реализовано с помощью авторской многопользовательской веб-игры «Душнота по понятиям», в которой участник письменно фиксировал собственное утверждение по причинно-следственному шаблону (формулируя некоторое понятие или взаимосвязь), оценивал ясность своего замысла по шкале, отвечал на «глупые» вопросы других участников (простые, неожиданные вопросы, проверяющие границы понимания) и затем оценивал ясность повторно — текущую (насколько ясно сейчас) и ретроспективную (насколько ясно было в начале до обсуждения), что позволяет зафиксировать сдвиг в самооценке. Данные собраны на двух отраслевых конференциях, CodeFest и PeopleSense, с участием специалистов в области IT и дизайна: полные наборы самооценок получены от 88 и 41 участника, дополнительно собрано 226 ответов в сравнительном условии, где те же вопросы участник применял к собственному тексту без внешнего собеседника, что позволило разделить эффект социального взаимодействия и эффект саморефлексии. Обработка выполнена парным t-критерием Стьюдента с расчётом d_z Коэна для оценки величины эффекта, критерием хи-квадрат для анализа распределений ответов и критерием Джонкхира – Терпстры для выявления тенденций в порядковых данных. Показано, что после обсуждения участники оценивают своё прежнее понимание ниже, чем в момент написания текста (+0,43 и +0,68 балла, $p = 0,003$), тогда как текущая оценка растёт; расхождение двух суждений (текущая оценка и ретроспективная оценка прежнего понимания) оказалось наиболее выраженным эффектом исследования, указывающим на то, что диалог с другими участниками разоблачает иллюзию понимания. При этом ни

количество вопросов, ни их субъективная полезность, ни число спрашивающих не связаны с величиной пересмотра, а в условии без собеседника сдвига не происходит вовсе, что свидетельствует о ключевой роли социального взаимодействия для запуска метакогнитивного пересмотра. Сделан вывод, что значение имеет не сам вопрос (как стимул для размышления), а внешняя позиция, из которой он поступает (Другой как носитель иной перспективы, выводящий из монологического мышления), и обязанность отвечать вслух, которая заставляет артикулировать и проверять границы собственного понимания, что делает игру эффективным инструментом самодиагностики и обучения взрослых.

Для цитирования в научных исследованиях

Агамалиев Р., Дорофеев М. Иллюзия объяснительной глубины на материале собственного замысла: самодиагностика в группе с помощью веб-игры // Психология. Историко-критические обзоры и современные исследования. 2026. Т. 15. № 6А. С. 76-99. DOI: 10.34670/AR.2026.29.68.024

Ключевые слова

Иллюзия понимания, иллюзия объяснительной глубины, самооценка понимания, ретроспективная оценка, метапознание, беглость суждения, внешний собеседник, веб-игра, обучение взрослых.

Введение

В начале 2026 года при подготовке к написанию настоящего исследования нами проведен небольшой неформальный эксперимент средствами интернета. Авторами настоящего исследования была написана статья [Агамалиев, 2026], в которой в достаточной для визуального представления степени была описана история изобретения первой цифровой камеры Стивеном Сассоном. Далее участникам/читателям статьи предлагалось сформулировать письменно и отправить через форму на сайте сообщение о том, что они поняли после прочтения. Было получено 108 ответов, из которых лишь в трех респонденты написали, что не знают, что произошло, и не понимают.

Остальные сто пять респондентов не проверяли сделанных авторами неформального эксперимента утверждений. Вместо проверки респонденты второй группы просто перефразировали предложенное нами объяснение, которое было ложным (о чём было написано в статье ниже). При проектировании финальной экспериментальной части статьи исследователи ложное объяснение заменили фактологическим и подчеркнули, что даже знание фактов не позволяет сформулировать исчерпывающего объяснения того, что произошло в истории изобретения цифровой камеры Стивеном Сассоном.

В нашем неформальном эксперименте участники располагали лишь информацией, предложенной нами, и не испытывали потребности проверить написанное. По мнению Розенблита и Кейла, в случае с нашей экспериментальной статьей знание того “что” произошло было принято респондентами за понимание того “как” и “почему” произошло [Rozenblit, Keil, 2002]. Это явление нами описывается как “иллюзия понимания”, и в данной работе мы будем оперировать именно этим понятием, суть которого в определенной степени совпадает с иллюзией объяснительной глубины [Rozenblit, Keil, 2002].

Розенблит и Кейл, изучая иллюзию объяснительной глубины, исследовали то, насколько хорошо респонденты понимают работу технических устройств, а в нашей работе мы рассматриваем иллюзию понимания при объяснении собственных решений в повседневной профессиональной деятельности.

Например, в случаях, когда кто-то говорит, что команда или ученик недостаточно замотивирован, чтобы выполнить ту или иную задачу. Большинству тех, кто услышит подобное утверждение, может показаться, что они понимают, что происходит, однако для того чтобы развеять иллюзию достаточно задать два вопроса:

- В чем выражается (Как выглядит) недостаточность мотивации?
- Как работает мотивация?

Два вопроса выше представляют собой категорию «глупых» вопросов (классификация авторская, полный список вопросов смотри в приложении №1 настоящего исследования). Говорящему может показаться, что он понимает суть описываемых им явлений, но при попытке ответить на «глупые» вопросы из приложения №1 он начинает явно замечать недостаточность знаний и понимания фактов, явлений и событий. Аналогичный эффект возникает, когда в отношении утверждения «Кто-то недостаточно замотивирован для выполнения какой-то задачи» просят предоставить обоснование в виде элементарной причинно-следственной нотации: Если кто-то недостаточно замотивирован, то он/она не выполняет какое-то действие, потому что ...?

Объяснение того, как что-то функционирует, или попытка построения логического умозаключения, по нашему мнению, помогает преодолеть иллюзию понимания, снизить чрезмерную уверенность в знании процессов объясняемого события, явления или факта, а также помогает индивиду преодолеть проблему, с которой он сталкивается при объяснении чего-либо. Проблему, которая заключается в том, что объяснение имеет сложную иерархическую структуру, при которой один элемент объяснения объясняет предыдущий, а финальное состояние объяснения никем не формализовано.

Иными словами, у объяснения отсутствуют критерии, при которых оно становится исчерпывающим, то есть у объяснения отсутствует состояние, при котором можно остановиться. Отсюда следует проблема: раз внешний критерий исчерпаемости отсутствует, значит, тот, кто объясняет, использует внутренний, которым становится легкость, с которой феномен объясняется. Шварцем в работе о метакогнитивных переживаниях этот феномен описан как «беглость как источник суждения» [Schwarz, 2004].

В данной работе мы утверждаем, что знание, которое обычно является основой иллюзии понимания, извлекается из памяти, а человеческая память не является надежным источником истины [Bartlett, 2003]. Джордан К. считает, что семантический контекст, в котором рассматривается явление, факт или событие, поддерживает процесс воспроизведения, который в условиях недостатка времени и обилия внешних раздражающих факторов создает исключительно иллюзию понимания фактов, событий или явлений. Иллюзию такого уровня и силы, при которой просмотр видеоинструкции сложнейшего процесса приземления самолета малой авиации создает у человека незнакомого с авиацией уверенность в том, что если потребуется он сможет самостоятельно посадить самолет [Jordan et al., 2022]. К аналогичному выводу пришел русский логик и философ С.И. Поварнин, в своей работе «Как читать книги» [Поварнин, 1924] он выдвинул тезис, что представить что-то приравнивается к пониманию. То есть феномен, при котором визуальная составляющая замещает логическое обоснование того, почему что-то функционирует так, как функционирует, существовал задолго до появления

цифровых технологий.

Информация, попадающая в поле внимания индивида (слова, обрывки изображений, звуки), собирает определенный семантический контекст, который в свою очередь запускает воображение индивида, которым компенсируется недостаток опыта. Чем меньше личного опыта и знаний, тем больше доля воображения и тем выше субъективная уверенность, тем сильнее иллюзия понимания [Kruger, Dunning, 1999].

Правдоподобным выглядит предположение о том, что эффект иллюзии понимания возник в связи с развитием цифровой среды, однако нами не найдено подтверждения тому, что иллюзия понимания является артефактом цифровой эпохи. Тем не менее, несмотря на то что цифровая среда не создает описанного выше механизма, цифровые средства существенно удешевляют возникновение эффекта иллюзии понимания. Мгновенный доступ к внешнему источнику информации распределяет познавательную нагрузку, при которой часть функций передается внешнему носителю информации; Вегнером этот эффект описывается через транзактивную память [Wegner, 1987], а частным случаем перераспределения выступает «эффект Гугла» [Sparrow, Liu, Wegner, 2011]: информация, доступная по запросу, воспроизводится хуже, чем путь к ней. Следствием иного порядка принято считать, факт, при котором использование интернета как сверхлёгкого способа получения информации приводит к тому, что человек стирает границы между своими собственными возможностями и возможностями электронных устройств, приписывая себе их сверхспособности [Hamilton, Yao, 2018].

В рамках данного исследования имеет значение не сам эффект ослабления памяти, а то, что внешние источники информации и цифровая среда стали частью семантического контекста, с помощью которого индивид формирует личностное понимание.

Доступность информации подменила личный опыт иллюзией личного опыта: готовое объяснение, полученное из внешней (цифровой) среды, создает эффект легкости, чего невозможно получить, если формулировать объяснение самостоятельными умственными усилиями. Чем дешевле достается объяснение, тем ярче у индивида проявляется иллюзия понимания.

Мы осознаем, что самостоятельное объяснение феномена, явления или факта как инструмент преодоления иллюзии понимания несет в себе ряд рисков, наиболее существенный из которых заключается в том, что проверять приходится тем же инструментом, которым иллюзия создается. При отсутствии критерия исчерпаемости объяснения индивид прибегает к наиболее доступному – легкости, с которой объяснение приходит на ум. Соответственно, адекватность объяснения оценивается по беглости, которая имитирует полноту [Alter, Oppenheimer, 2009], замыкая тем самым петлю: чем свободнее человек говорит о явлении, факте или событии, тем убедительнее ему кажется, что он его понимает.

Разрыв петли самостоятельно возможен, но исключительно требователен к умению задавать вопросы к собственным убеждениям и умению оценивать логичность ответа. Умение спрашивать и оценивать состоятельность собственного ответа формируется годами и требует специальной подготовки. Значительно более доступным решением оказалось привлечение другого собеседника, партнера, задающего «глупые» вопросы. Внешний собеседник обладает рядом преимуществ перед тем, кто объясняет самому себе: он не разделяет семантический контекст, в котором сделано утверждение, возник факт или явление, и не имеет доступа к чувству легкости, сопровождающему утверждение автора. Ему доступны только сказанные или написанные слова, и оценивать понимание он вынужден исключительно по утверждению.

Кроме того, позиция защищающего утверждение и позиция оспаривающего его без

длительной тренировки не могут быть воспроизведены одним человеком одновременно. Метод последовательного опровержения, известный со времен сократического эленхоса [Farnsworth, 2021], предполагает наличие оппозиции, в этом отношении диалог располагает ресурсом, которого лишена письменная саморефлексия. Именно поэтому в настоящем исследовании интервенция построена не как самостоятельное объяснение [Fernbach et al., 2013], к которому прибегают Розенблит и Кейл [Rozenblit, Keil, 2002], а как серия ответов на “глупые” вопросы, задаваемые другими участниками (партнерами).

Настоящее исследование проверяет три следствия изложенного выше:

Первое. Если объясняющий оценивает полноту собственного понимания по беглости, а внешний собеседник такого доступа лишен и располагает только словами, то оценки, выставленные докладчику участниками, задающими вопросы, окажутся ниже его собственной, при этом расхождение сохранится и после обсуждения.

Второе. Мы считаем, что если письменная саморефлексия в основе имеет непригодный критерий оценки понимания, а диалог вводит внешнюю (независимую) позицию, то задание на письменную рефлексия не изменит самооценку понимания, тогда как серия вопросов, заданных другим участником, изменит.

Третье. Если “глупые вопросы” разрушают неосознанную индивидом иллюзию понимания при первоначальной оценке, то участник, которому после ответов на “глупые” вопросы предложат оценить, каким его понимание было в самом начале, поставит оценку ниже первоначальной.

Измеряемой величиной во всех трех следствиях выступает самооценка глубины понимания собственного текста по шкале Ликерта от 1 до 7, собираемая до обсуждения, непосредственно после обсуждения и ретроспективно.

Ретроспективная оценка требует оговорки, так как запрашивается вопросом, прямо указывающим на возможность обнаружения пробелов, и потому подвержена как влиянию наводящей формулировки, так и ретроспективному искажению сути. Мы вернемся к этому при обсуждении ограничений настоящего исследования.

Материалы и методы исследования

Дизайн исследования предполагал групповое измерение с независимым воспроизведением на двух площадках. Участие в эксперименте было добровольным и выглядело для участника, как участие в обучающем мастер-классе. Каждый участник последовательно выступал в двух ролях: автора собственного текста (докладчик) и задающего вопросы автору другого текста (душила). Самооценка глубины понимания собственного текста измерялась трижды: до обсуждения, сразу после обсуждения и ретроспективно, применительно к моменту написания текста.

В качестве дополнительной меры измерения на первой площадке было реализовано сравнительное условие, в котором тот же инструментальный вопрос предьявлялся участникам для самостоятельного использования в отношении собственного текста, без диалога с другим человеком.

Сбор данных осуществлялся на двух профессиональных конференциях: CodeFest (31 мая 2026 г.) и PeopleSense (5 июня 2026 г.). Обе площадки являются открытыми отраслевыми конференциями; участие в исследовательском эксперименте было добровольным и не оказывало никакого влияния на посещение других докладов и мастер-классов программы

конференции.

Экспериментальное вмешательство реализовано с помощью авторской многопользовательской веб-игры «Душнота по понятиям» (nitpicker.mnogodelal.ru). Регистрация ограничивалась вводом имени или псевдонима, сбор иных персональных данных не осуществлялся. Все действия участников – тексты, вопросы, оценки, этапы смены ролей – автоматически фиксировались в базе данных SQL с отметками времени.

Так как площадок было две, нам кажется логичным детализировать данные участников по каждой из площадок отдельно.

CodeFest. 137 человек прошли регистрацию в игре, двадцать один из которых не оставил ни одной оценки – это прерванные или повторные подключения. Написали свой текст и дали ему начальную оценку 116 участников – то есть фактически именно они вошли в эксперимент, – и далее мы используем это число как основание выборки. Полный ответ (оценка понятности собственного текста до, после, ретроспективно) получен от 88 участников, что составляет 75,9% от числа вошедших в процедуру.

На CodeFest получено 211 оценок понятности текста докладчика, в которых участвовали 118 из 137 зарегистрированных. Из них 111 входят в число 116 вошедших в процедуру, а оставшиеся семь принадлежат участникам, подключившимся в роли душнили минуя первую фазу (о фазах далее) и потому не имеющим начальной самооценки.

PeopleSense. 68 человек зарегистрировались в игре. Текст написали 56 участников, все три оценки понятности собственного текста (до, после, ретроспективно) получены от 41 участника, что составляет 73,2%.

В исследовании использовался материал, созданный (написанный) самими участниками непосредственно перед началом эксперимента, что принципиально отличает наш дизайн исследования от исследования Розенблита и Кейла, в котором объектом объяснения выступает внешний к участнику объект.

На CodeFest участникам было предложено написать обоснование личного участия в конференции. На PeopleSense участники формулировали свои опасения от следования практикам гигиены умственного труда, описанных Максимом Дорофеевым в научно-популярной книге “Путь джедая”, – отключение уведомлений, провести утро без смартфона, удаление приложений социальных сетей и т.д. Средняя длина текста составила 408 символов (от 265 до 725). Текст строился по шаблону, в основе которого лежала базовая причинно-следственная конструкция: если ..., то..., потому что... .

Примеры текстов участников приводятся дословно, с сохранением авторской орфографии и пунктуации. Имена не указываются: при регистрации участники вводили произвольное самоназвание, и оно не несёт исследовательской информации. В скобках после каждого текста даны три самооценки этого же участника в порядке сбора – начальная, после обсуждения и ретроспективная.

На CodeFest тексты строились по четырёхчленной нотации и содержательно представляли собой обоснование собственного участия в конференции.

Если я приеду на CodeFest, то познакомлюсь с людьми которые создают новые технологии в ИТ, потому что на кодфесте более 200 спикеров экспертов

А это нужно, чтобы я при решении задачи имел широкий кругозор и мог выбрать подходящий под задачу инструмент, технологию или подход,

так как при общении с экспертами я расширяю свои познания в ит и связанных областях. (5 – 6 – 3: после обсуждения участник оценил своё прежнее понимание на две позиции ниже.)

Если я побываю на CodeFest, то получу интересные инсайты как минимум в одной из интересующих меня областях, а именно управление командами, применение ИИ в разработке, техники мышления и отдыха, потому что здесь обычно хороший подбор спикеров на интересующие меня темы. А это нужно, чтобы расширять кругозор, так как кругозор помогает мыслить шире и, в итоге, иметь больший диапазон решений для встречающихся на работе и в личной жизни задач. (5 – 5 – 5: ни одна из оценок не изменилась.)

Если я буду слушать лекции про фронтенд на кодфест, то изучу новые технологии фронта, потому что буду внимательно вникать. А это нужно, чтобы внедрить новые решения в работе, так как ускорение разработки повысит мою продуктивность и ценность как профессионала. Таким образом компания сможет заработать больше денег и я смогу получить премию, что улучшит моё материальное положение. (5 – 5 – 6: ретроспективная оценка выше начальной, то есть пересмотр произошёл в обратную сторону.)

На PeopleSense нотация была трёхчленной, а предметом высказывания выступало опасение получить негативное последствие от следования практике – то есть одно из возможных оснований для отказа от применения практики в своей повседневной жизни.. Такой формат высказывания продиктован педагогическим взглядом авторов на обучение взрослых людей «мягким» навыкам: от изменения поведения (именно его авторы считают целью обучения) ученика удерживает не недостаток знания, а сопротивление, часто иррациональное.

Угловые скобки в двух из трёх примеров – остаток предъявленного участникам шаблона, который они не удалили.

ЕСЛИ я отключу уведомления о новых сообщениях

ТО я упущу что-то важное для меня, какую-то новость, упущу встречу, событие

ПОТОМУ ЧТО что-то важное пройдет мимо меня, я потеряю контроль

(5 – 7 – 5.)

ЕСЛИ <Я регулярно закрываю все вкладки в браузере>,

ТО <потом я потеряю важную информацию и не смогу её найти>,

ПОТОМУ ЧТО <вкладки постоянно мне об этой информации будут напоминать. И можно будет легче к этому вернуться.>

(7 – 4 – 4: случай, в котором обсуждение снизило не только ретроспективную, но и текущую оценку.)

ЕСЛИ буду делать только записанное в список задач,

ТО то я буду больше времени тратить на запись и ранжирования мелких и средних задач,

ПОТОМУ ЧТО много задач которые делаются быстро проще сделать, могут осмысливаться долго, или будут в последствии не нужны

(4 – 6 – 4.)

Процедура

Дизайн эксперимента выглядел следующим образом:

Участникам было предложено объединиться в малые группы по 2-4- человека,

Работа в группах проходила в три раунда, каждый раунд состоял из пяти последовательных фаз.

Между раундами докладчиком становился следующий участник группы, все остальные участники группы становились «душилами».

. На обеих площадках дизайн был одинаков.

Фаза 1. Текст и начальная самооценка. Участник фиксировал свою мысль в текстовой форме, после чего отвечал на вопрос: «Насколько вам ЯСЕН этот план на данный момент?»

Ответ фиксировался как [initial_clarity].

Фаза 2. Распределение ролей. Участник, выбравший роль докладчика, получал случайный шестизначный код; один или два других участника группы, ставшие «душилами», вводили этот код для подключения к сессии докладчика. Число «душилов» на одного докладчика регламентировалось рекомендацией «от одного до двух», фактически распределение составило: на CodeFest – от 1 до 6, на PeopleSense – от 1 до 3. На CodeFest в среднем на одного докладчика приходилось 2,27 душилов, на PeopleSense – 1,64.

Фаза 3. Вопросы и оценка по тексту. По сценарию эксперимента, после того как душиловы подключались к сессии докладчика, им предъявлялся его текст, в котором необходимо было выделять отдельные слова и формулировать к ним вопросы. Формулирование поддерживалось набором шаблонов, привязанных к частям речи и типам конструкций: существительное, глагол, прилагательное, наречие, местоимение, число или дата, абстрактное понятие, метафора или образ, логическая связка, конструкция «все/всегда», модальная конструкция «надо/должен/нельзя».

Завершив раунд формулирования вопросов к словам и понятиям из утверждения докладчика, «душила» должен был ответить на вопрос: «Судя по тексту докладчика, насколько он, по вашему мнению, понимает свой план?» Ответ фиксировался как [nitpicker_pre_clarity] – отдельно для каждого «душиловы».

Всего на CodeFest было задано 863 вопросов, из которых 604 получили ответ; на PeopleSense задали 332 вопроса и получили 246 ответов.

Фаза 4. Обсуждение. Докладчик видел к каким словам в его тексте у «душилов» возникли вопросы, но не видел текст вопросов, пока не выберет конкретного слова. Как только докладчик выбирал слово, вопрос к которому хотел обсудить, у всех участников в приложении показывался текст вопроса. Докладчик устно отвечал на вопрос и оценивал его полезность по семибалльной шкале, а «душиловы» оценивали полноту ответа докладчика.

Фаза 5. Итоговые оценки. Завершающим этапом, после ответов на все полученные вопросы или на те из них, на которые он успел ответить до смены ролей, докладчик отвечал на два вопроса: а) «Насколько вам ясен этот план СЕЙЧАС, после обсуждения?» [current_clarity], б) «Вспомните момент, когда вы писали свой текст. Как на самом деле вам стоило бы оценить уровень ясности ТОГДА?». Второй вопрос задавался с небольшим уточнением, что «сейчас, после вопросов «душиловы», вы могли обнаружить проблемы в своём понимании, которых не замечали ранее» [retrospective_clarity].

«Душиловы» отвечали на вопрос: «А теперь, после обсуждения, как ты думаешь, насколько хорошо докладчик НА САМОМ ДЕЛЕ понимает свой план?» [nitpicker_post_clarity].

Все пять измеряемых показателей собирались по шкале Ликерта, где 1 означало «совсем не ясен, много неопределённости», а 7 – «полностью ясен, всё конкретно и определённо».

Таблица 1 - Измеряемые показатели в порядке сбора

Показатель	Кто оценивает	Кого оценивает	Момент измерения
initial_clarity	докладчик	себя	после написания текста, до обсуждения
nitpicker_pre_clarity	душила	докладчика	после прочтения текста и формулирования вопросов, до обсуждения
current_clarity	докладчик	себя	сразу после обсуждения
retrospective_clarity	докладчик	себя в момент написания текста	сразу после обсуждения
nitpicker_post_clarity	душила	докладчика	сразу после обсуждения

Оценки, полученные докладчиком от нескольких «душиль», усреднялись, и в анализ входило одно усреднённое значение на докладчика. Полезность каждого вопроса оценивалась докладчиком также по семибалльной шкале.

На основе пяти измеряемых показателей рассчитывались четыре производные величины (таблица 2). Оценки, полученные докладчиком от нескольких «душиль», предварительно усреднялись, так что в расчёт входило одно значение внешней оценки на докладчика.

Таблица 2 - Производные показатели

Показатель	Формула	Что отражает
IOED_speaker	$\text{initial_clarity} - \text{retrospective_clarity}$	Величину ретроспективного пересмотра: насколько ниже участник оценивает своё прежнее понимание после обсуждения
clarity_gain	$\text{current_clarity} - \text{retrospective_clarity}$	Субъективный прирост понимания относительно пересмотренной оценки прошлого состояния
pre_bias	$\text{nitpicker_pre_clarity} - \text{initial_clarity}$	Расхождение внешней оценки по тексту с самооценкой автора до обсуждения
post_bias	$\text{nitpicker_post_clarity} - \text{current_clarity}$	То же расхождение после обсуждения

Знак производных показателей интерпретируется следующим образом. Положительное значение [IOED_speaker] означает, что после обсуждения участник оценил своё прежнее понимание ниже, чем оценивал его в момент написания текста, то есть обнаружил у себя иллюзию понимания. Отрицательные значения [pre_bias] и [post_bias] означают, что внешние наблюдатели оценивают понимание докладчика ниже, чем он оценивает себя сам.

Показатель [clarity_gain] требует оговорки. Обе входящие в него оценки собираются в один и тот же момент – сразу после обсуждения, – и потому он отражает не изменение понимания во времени, а различие между двумя суждениями участника: о своём состоянии сейчас и о своём состоянии тогда. Мы используем его как описательную характеристику расхождения этих двух суждений, а не как меру научения.

Следует отметить различие между площадками. На CodeFest мастер-классу предшествовал доклад, в ходе которого те же шаблоны вопросов (в сокращённом наборе) предъявлялись аудитории со сцены. Процедура на докладе была практически идентичной той, что была на мастер-классе: участники письменно фиксировали план (обоснование своего участия на конференции), а затем шел рассказ о шаблонах вопросов. После рассказа о каждом шаблоне была пауза на 2 минуты, чтобы участники отмечали в своём тексте слова, к которым возникают вопросы, в соответствии с этим шаблоном. Участники доклада не задавали вопросы к тексту другого человека, только к своему. В ходе работы участник отвечал на вопрос о глубине понимания до рассказа о шаблонах и после, оценивая её по той же семибалльной шкале.

Ответы на докладе CodeFest собирались через анонимную форму: получено 226 полных ответов.

Принципиальное отличие условий доклада от условий мастер-класса заключается в источнике вопроса. При одинаковом инструменте во время доклада участник применял его к собственному тексту самостоятельно, без внешнего собеседника. Подобный дизайн позволяет нам изолировать действие вопросов от действия позиции спрашивающего.

Существенное ограничение, к которому мы ещё вернёмся, заключается в том, что форма ответов была анонимной, без индивидуальной связки волн (доклад, мастер-класс). В связи с

этим сравнение возможно исключительно на уровне групповых распределений. Кроме того, с очень высокой вероятностью заметная часть участников мастер-класса опробовала игру в рамках доклада, то есть условия не являются независимыми, а образуют последовательность. Порядок условий не менялся.

Групповое сравнение проводилось парным t-критерием Стьюдента, в качестве меры величины эффекта рассчитывался d_z Коэна. Сравнение распределений сдвига самооценки между условиями проводилось критерием хи-квадрат с расчётом V Крамера. Связь количественных показателей оценивалась коэффициентом корреляции Пирсона. Критический уровень значимости принят равным 0,05. Расчёты выполнены в XLSTAT 2019.4.2 и независимо проверены средствами Python.

Результаты и обсуждение

Основную выборку составили участники, ответившие на все три вопроса самооценки: 88 участников мастер-класса CodeFest и 41 – PeopleSense. Отсев на промежуточных фазах не был связан с исходной самооценкой понимания написанного текста.

На CodeFest дошедшие до ретроспективной оценки участники оценили ясность своего плана в среднем в 5,16 балла ($n = 88$), выбывшие – в 4,86 ($n = 28$). Различие не достигает уровня статистической значимости, $t = 0,89$, $p = 0,378$, $d = 0,23$. На PeopleSense соответствующие значения составили 5,54 ($n = 41$) и 5,48 ($n = 21$), $t = 0,19$, $p = 0,851$, $d = 0,05$. Иные причины отсева проверить по другим показателям невозможно, поскольку в дизайне эксперимента подобный показатель отсутствовал.

Данные, полученные с обеих площадок, подтвердили нашу гипотезу о том, что самооценка ясности собственного текста после обсуждения должна возрасти, а ретроспективная оценка первоначального написанного текста снизиться в сравнении с оценкой в момент написания.

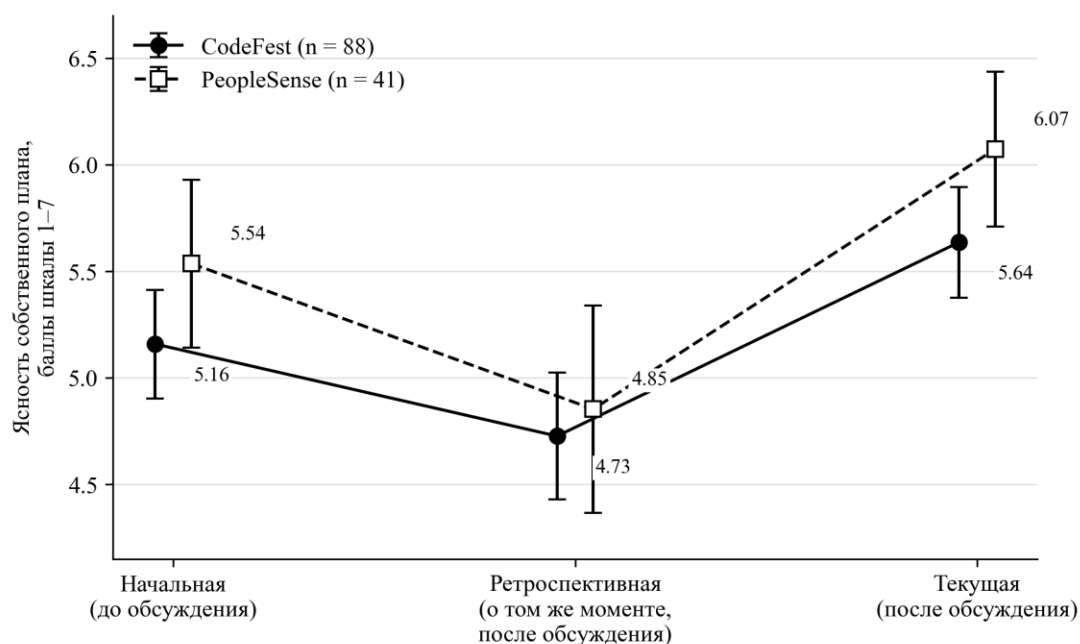


Рисунок 1 - Самооценка ясности собственного текста в трёх замерах на двух площадках, средние значения и 95% доверительные интервалы

Таблица 3 - Самооценка ясности собственного текста и её изменения

Показатели	CodeFest (n = 88)	PeopleSense (n = 41)
Средние значения, М (SD)		
Начальная, до обсуждения	5,16 (1,20)	5,54 (1,25)
Текущая, после обсуждения	5,64 (1,22)	6,07 (1,15)
Ретроспективная	4,73 (1,40)	4,85 (1,54)
Ретроспективный пересмотр, IOED_speaker		
Δ, баллы [95% ДИ]	+0,43 [0,15; 0,71]	+0,68 [0,25; 1,12]
Критерий	t(87) = 3,03; p = 0,003	t(40) = 3,20; p = 0,003
d_z Коэна	0,32	0,50
Прирост текущей оценки над начальной		
Δ, баллы [95% ДИ]	+0,48 [0,22; 0,74]	+0,54 [0,16; 0,92]
Критерий	t(87) = 3,69; p < 0,001	t(40) = 2,85; p = 0,007
d_z Коэна	0,39	0,45

Примечание. Положительное значение IOED_speaker означает, что после обсуждения участник оценил своё прежнее понимание ниже, чем оценивал его в момент написания текста.

В соответствии с полученными данными, ретроспективная оценка и оценка понимания после обсуждения сдвигаются в противоположных направлениях. Участник после обсуждения сообщает, что понимает свой текст лучше, чем понимал до обсуждения, и в то же время указывает на то, что понимал хуже, чем ему до этого казалось. Мы рассчитываем сдвиг с помощью показателя clarity_gain:

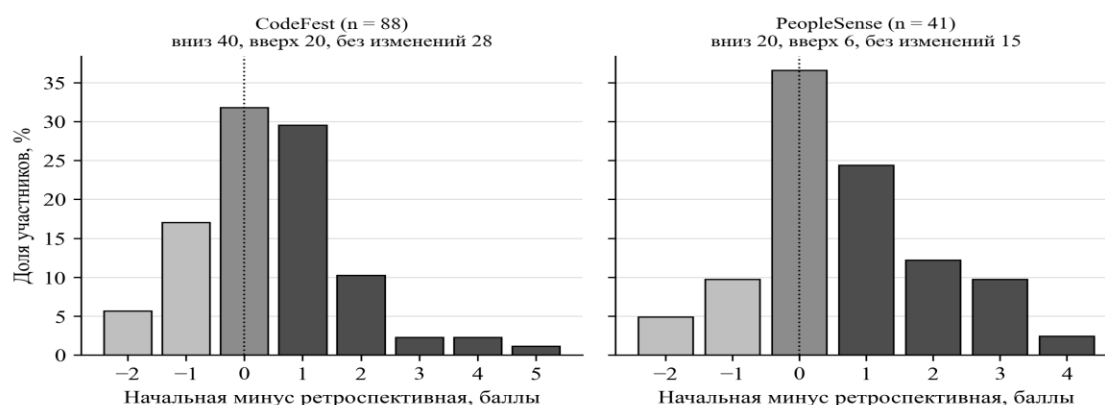
Таблица 4 - Эконометрические показатели исследования

clarity_gain	CodeFest (n = 88)	PeopleSense (n = 41)
Δ, баллы [95% ДИ]	+0,91 [0,67; 1,15]	+1,22 [0,80; 1,63]
Критерий	t(87) = 7,50; p < 0,001	t(40) = 5,94; p < 0,001
d_z Коэна	0,80	0,93

Обе оценки (после обсуждения и ретроспективная) получены в один и тот же момент, а потому этот показатель описывает расхождение суждений, а не изменение понимания.

Структура пересмотра первоначальной оценки

Природа ретроспективной оценки вынуждает нас выделить то обстоятельство, что пересмотр в сторону уменьшения не был общим свойством выборки.

**Рисунок 2 - Распределение ретроспективного пересмотра начальной оценки (начальная минус ретроспективная) на двух площадках**

На CodeFest ретроспективную оценку относительно первоначальной понизили 40 участников из 88 (45%), повысили 20 (22%) и не изменили 28 (32%). На PeopleSense понизили 20 из 41 (49%), повысили 6 (15%) и не изменили 15 (37%).

Такое же распределение выявлено и для текущей оценки. Хотя в среднем она выросла, у 15 участников CodeFest (17%) и 7 участников PeopleSense (17%) она оказалась ниже начальной, то есть примерно у каждого шестого участника обсуждение понимания не добавило, а, наоборот, убавило.

Распределение самооценок по баллам уточняет характер сдвига:

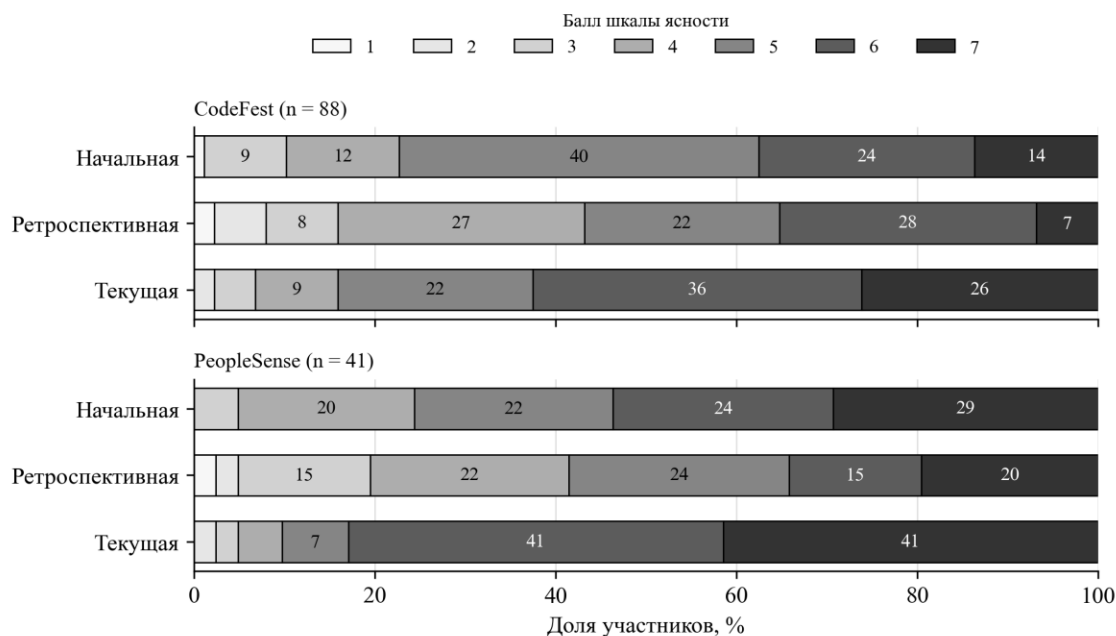


Рисунок 3 - Распределение самооценок по баллам шкалы ясности в трёх замерах на двух площадках

Ретроспективная оценка отличается от начальной не отрицательным смещением всей выборки, а перераспределением части выборки из зоны средних баллов (середины шкалы) в нижние баллы.

Внешняя оценка ясности ниже, чем самооценка, но после совместного обсуждения расхождение с первоначальной оценкой уменьшилось.

На CodeFest до обсуждения различие не достигает значимости: 5,16 против 4,89, $\Delta = 0,26$, $t(87) = 1,52$, $p = 0,131$; однако после обсуждения оно становится значимым: 5,64 против 5,27, $\Delta = 0,37$, $t(87) = 2,25$, $p = 0,027$, $d_z = 0,24$.

На PeopleSense различие значимо при обоих замерах. До обсуждения самооценка составила 5,53 балла против 4,85 у душили, $\Delta = 0,68$, $t(40) = 3,09$, $p = 0,004$, $d_z = 0,49$, а после обсуждения 6,05 против 5,37, $\Delta = 0,68$, $t(40) = 3,17$, $p = 0,003$, $d_z = 0,51$.

Существенным фактором в защиту нашей гипотезы мы считаем, что после обсуждения оценки респондентов на обеих площадках выросли, как собственная, так и внешняя. Душили, выслушавшие докладчика, стали оценивать его понимание выше, а не ниже, при этом разрыв между самооценкой и оценкой со стороны сохранился.

Сравнение с условием без собеседника (CodeFest доклад)

В условиях, когда участник применял вопросы к себе без участия «душиль», средняя оценка понимания практически не изменилась: 5,11 (SD = 1,44) до предъявления карточек и 5,19 (SD = 1,63) после, $\Delta = +0,08$, 95% ДИ [-0,10; 0,27], $t = 0,88$, $p = 0,380$, $d_z = 0,06$ ($n = 226$).

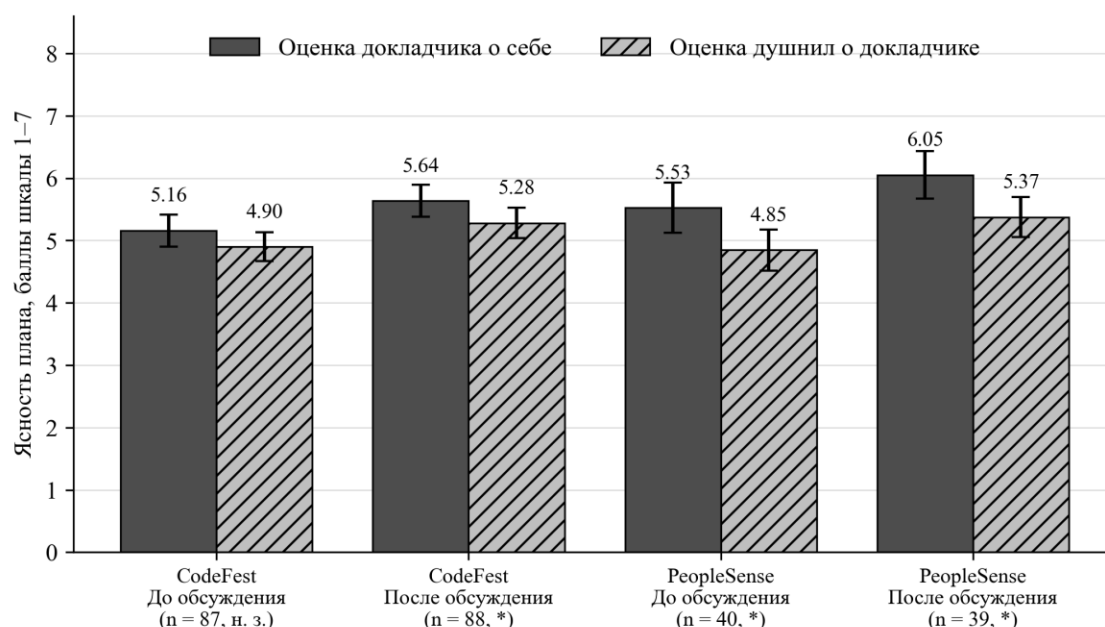


Рисунок 4 - Самооценка докладчика и оценка душиль до и после обсуждения, средние значения и 95% доверительные интервалы

Мы считаем, что отсутствие видимого результата при сравнительном условии не означает, что инструмент «не работает». Вопросы к собственному тексту возникли у 199 респондентов, то есть у 88% выборки, в среднем к 26,5% слов; лишь 12% не отметили ни одного слова в своих формулировках.

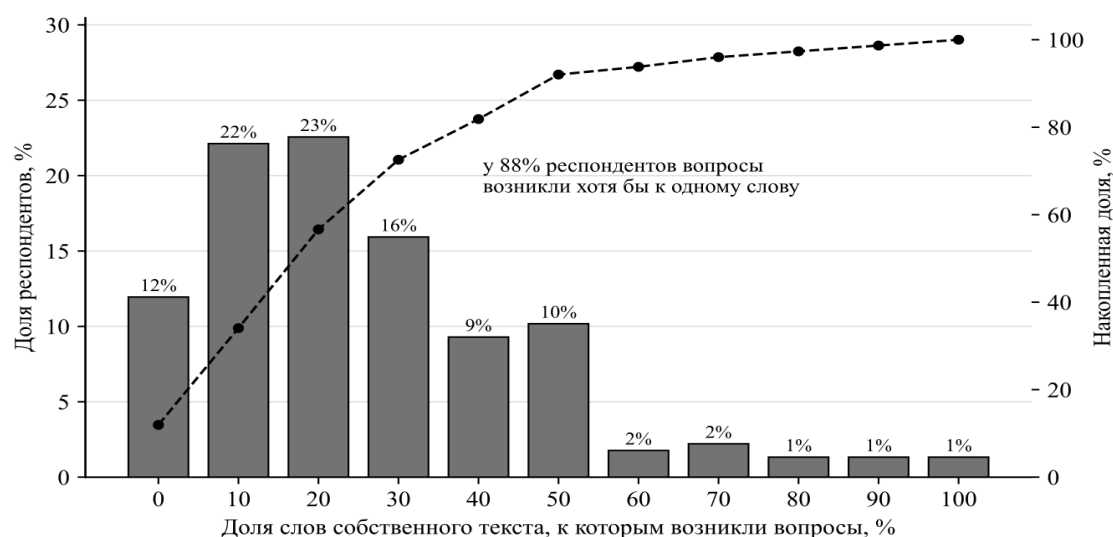


Рисунок 5 - Доля слов собственного текста, к которым у респондента возникли вопросы, и накопленная доля респондентов (условие без собеседника)

Инструмент решил задачу: участники обнаружили в своём тексте непонятные слова и формулировки, тем не менее самостоятельная работа с текстом среднюю оценку понимания не изменила. Мы предположили, что причина заключается в том, что, увидев вопросы, респонденты посчитали, что ответы на большинство из них они уже знают.

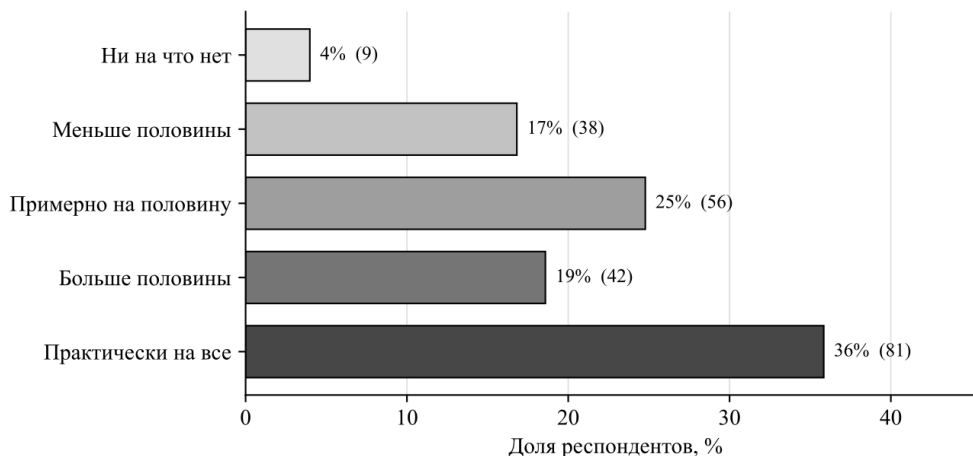


Рисунок 6 - Распределение ответов на вопрос «На какую часть вопросов у вас есть ответ» (условие без собеседника)

36% ответили, что у них есть ответы практически на все вопросы к собственному тексту, 19% – что больше чем на половину, и лишь 4% посчитали, что ответов нет ни на один вопрос. Само по себе полученное распределение не проверяет предположение, но указывает на тот факт, что участники уверены в наличии у них ответов, а также отметим, что распределение по шкале «от практически на все» до «ни на что» не показывает, связана ли эта уверенность со сдвигом оценки. Мы решили проверить наше предположение, используя массив полученных данных.

Ответ на вопрос относительно того, на какую часть вопросов у респондентов есть готовые ответы, представляет собой порядковую шкалу из пяти градаций, благодаря этому мы можем сопоставить её со сдвигом самооценки и проверить упорядоченную альтернативу, а именно: возрастает ли сдвиг по мере роста уверенности в наличии ответов. Для анализа мы использовали критерий Джонкхир – Терпстры.

Анализ подтверждает наличие тренда – данные демонстрируют монотонное возрастание: $J = 12\,423,5$ при $E(J) = 9\,522,5$, $z = 5,50$, $p < 0,001$, $\tau\text{-}b$ Кендалла = 0,30. Результат воспроизводится при перестановочной проверке: ни одна из 50 000 перестановок не дала значения, равного наблюдаемому или превышающего его ($p < 0,0001$). Лучшая из них достигла $J = 11\,545$ при наблюдаемом $J = 12\,423,5$. Эффект сохраняется и при исключении наименьшей градации, на которую пришлось всего 9 человек, $z = 4,90$, $p < 0,001$.

Средний сдвиг растёт монотонно по всем пяти градациям: $-1,11$ балла $[-2,09; -0,14]$ у тех, у кого нет ответа ни на один вопрос, $-0,45$ и $-0,32$ в двух средних градациях, $+0,31$ и $+0,63$ $[0,34; 0,91]$ у тех, у кого есть ответы на большую часть вопросов и практически на все. Вместе со средним значением изменяется и направление сдвига. В нижней градации самооценку понизили 56% респондентов и не повысил никто, а в верхней 5% понизили, а повысили 44%. Объединение градаций попарно даёт тот же результат: $-0,57$ балла у тех, у кого ответов меньше половины ($n = 47$), против $+0,52$ у тех, у кого их больше половины ($n = 123$), $\Delta = -1,09$ $[-1,61; -0,58]$, $t = -4,20$, $p < 0,001$, $d = -0,78$.

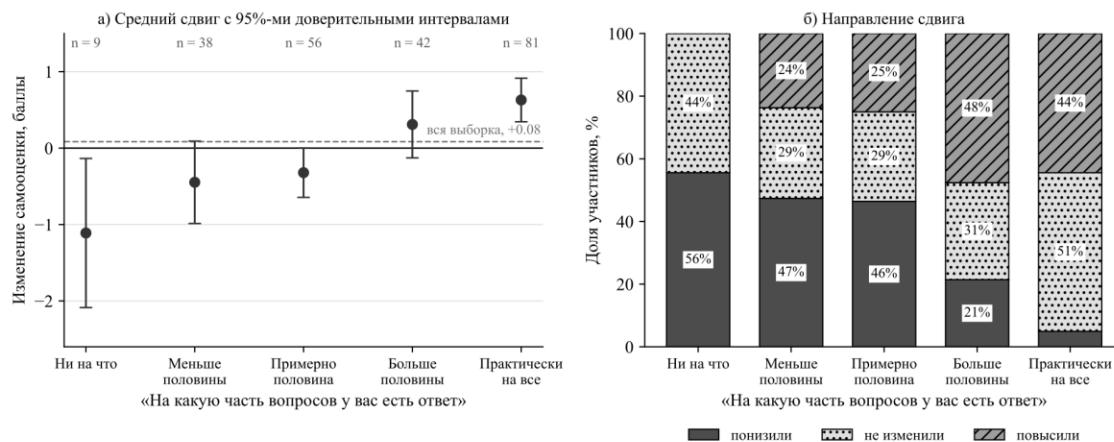


Рисунок 7 - Средний сдвиг самооценки и направление сдвига по градациям уверенности в наличии готовых ответов (условие без собеседника)

Из этого можно сделать следующий вывод: средний ноль в условиях без собеседника не означает отсутствие изменений. Он является результатом усреднения двух противоположно направленных подгрупп: признавшие отсутствие ответа понизили свою оценку на один балл, а посчитавшие, что ответы у них есть, повысили. Механизм самопроверки собственного текста включается, но у меньшинства – у двух нижних градаций выборки (21%).

Данный результат нельзя объяснить регрессией к среднему, поскольку данные идут в противоположную от регрессии сторону. Средняя первичная оценка по градациям составила 3,89, 4,18, 4,77, 5,48 и 5,73 балла при общем среднем 5,11. Регрессия при таких значениях поднимала бы нижние градации и опускала бы верхние, а наблюдается обратное. Контроль первичной оценки связь не ослабляет, а, наоборот, усиливает.

Мы считаем необходимым оговорить два обстоятельства. Первое: анализ эффекта самостоятельного оценивания имеет поисковый характер, так как не был предусмотрен при планировании и проведён после получения основных результатов. Второе: суждения “я понимаю свой текст” и “у меня есть ответы на эти вопросы” являются близкими по сути и могут рассматриваться респондентами как одно и то же. На второе обстоятельство указывает высокая связь градации с итоговой оценкой, $r = 0,69$.

Если принять во внимание всё вышесказанное, то можно сказать, что вопрос, заданный самому себе, является вопросом, на который, уже имеется готовый ответ.

Отдельно, в качестве меры дополнительной проверки предположения, что “если спрашивать самого себя, ответ уже имеется”, мы сравнили сдвиг самооценки при вопросах душнил к докладчику и при самостоятельном допросе.

Распределения сдвига самооценки в двух условиях значимо различаются: $\chi^2(4) = 19,71$, $p < 0,001$, $N = 315$, V Крамера = 0,25. После обсуждения самооценку повысили 60% респондентов против 35% в условиях доклада, понизили 17% против 27%, а не изменили 24% против 38%.

Обсуждение с внешним собеседником не столько предотвращает ухудшение самооценки, сколько добавляет умеренный прирост у тех респондентов, у которых она при самостоятельной оценке осталась бы без изменений.

Ограничения данного сравнения ранее оговаривались в разделе методологии. Связка между условиями групповая, а не внутрисубъектная, и порядок условий не менялся.

Всего в ходе игры «душилами» были заданы 863 вопроса на CodeFest и 332 на PeopleSense, из которых ответ получили 604 (70%) и 246 (74%) соответственно. На каждого докладчика

пришлось в среднем 8,99 вопроса на CodeFest (от 1 до 31) и 6,92 на PeopleSense (от 2 до 13).

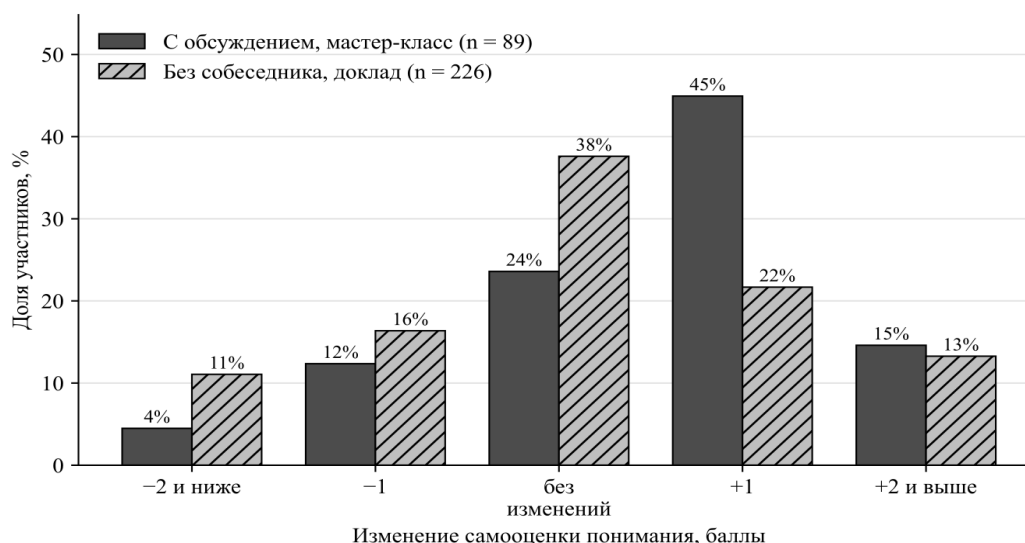


Рисунок 8 - Распределение сдвига самооценки понимания в двух условиях: с обсуждением (мастер-класс) и без собеседника (доклад).

Полезность вопросов докладчиками оценена: 5,83 из 7 на CodeFest ($SD = 1,40$, $n = 615$) и 5,81 на PeopleSense ($SD = 1,30$, $n = 244$). «Душлилы» полноту полученных ответов оценивали: 5,61 ($SD = 1,47$, $n = 1432$) и 5,59 ($SD = 1,21$, $n = 557$).

Три четверти использованных «душлилами» вопросов взяты из предъявленных им заранее заготовленных шаблонов. На CodeFest 217 из 863 (25%) и на PeopleSense 72 из 332 (22%) сформулированы самостоятельно, без шаблона. Наибольшее количество вопросов вызвали существительные (135 и 58), глаголы (119 и 51) и прилагательные (114 и 48).

Следует особо отметить, что ни один из показателей вмешательства (описанных разделом ранее) не предсказал величину ретроспективного сдвига. Корреляция IOED_speaker с числом полученных вопросов составила $r = 0,13$ ($p = 0,226$) на CodeFest и $r = 0,02$ ($p = 0,882$) на PeopleSense. Любая связь между количеством вопросов и оценкой их полезности не имеет статистической значимости, что заслуживает отдельного внимания.

Отсутствие связи между количеством и/или субъективным качеством вопросов и ретроспективным сдвигом оценки означает, что пересмотр происходит при наличии «душлилы» (внешнего спрашивающего), безотносительно к количеству и/или качеству вопросов. Получается, что сдвиг обеспечивает сам факт присутствия другого человека в позиции «душлилы» (спрашивающего), а не интенсивность и не воспринимаемое качество вопросов.

Таблица 4 - Сводка статистических сравнений

Сравнение	Условие	n	Δ [95% ДИ]	Критерий	p	Размер эффекта
Ретроспективный пересмотр: начальная – ретроспективная	CodeFest	88	+0,43 [0,15; 0,71]	$t(88) = 3,03$	0,003	$d_z = 0,32$
	PeopleSense	41	+0,68 [0,25; 1,12]	$t(40) = 3,20$	0,003	$d_z = 0,50$
Прирост самооценки: текущая – начальная	CodeFest	88	+0,48 [0,22; 0,74]	$t(88) = 3,69$	< 0,001	$d_z = 0,39$
	PeopleSense	41	+0,54 [0,16; 0,92]	$t(40) = 2,85$	0,007	$d_z = 0,45$
Расхождение двух суждений: текущая – ретроспективная	CodeFest	88	+0,91 [0,67; 1,15]	$t(88) = 7,50$	< 0,001	$d_z = 0,80$
	PeopleSense	41	+1,22 [0,81; 1,63]	$t(40) = 5,94$	< 0,001	$d_z = 0,93$

Сравнение	Условие	n	Δ [95% ДИ]	Критерий	p	Размер эффекта
Самооценка против оценки душили, до обсуждения	CodeFest	88	+0,26 [-0,08; 0,61]	t(87) = 1,52	0,131	d_z = 0,16
	PeopleSense	40	+0,68 [0,23; 1,12]	t(39) = 3,09	0,004	d_z = 0,49
Самооценка против оценки душили, после обсуждения	CodeFest	88	+0,37 [0,04; 0,69]	t(88) = 2,25	0,027	d_z = 0,24
	PeopleSense	39	+0,68 [0,25; 1,11]	t(38) = 3,17	0,003	d_z = 0,51
Изменение самооценки без собеседника: T1 – T0	Доклад	226	+0,08 [-0,10; 0,27]	t(225) = 0,88	0,380	d_z = 0,06
Тренд сдвига по уверенности в наличии ответов	Доклад	226	–	J = 12 423,5; z = 5,50	< 0,001	τ -b = 0,30
Сдвиг при ответах менее и более чем на половину вопросов	Доклад	47 123	-1,09 [-1,61; -0,58]	t(72,2) = -4,20	< 0,001	d = -0,78
Различие распределений сдвига между условиями	Мастер-класс против доклада	315	–	$\chi^2(4) = 19,71$	< 0,001	V = 0,25

Примечание. Для первых трёх сравнений положительное Δ означает, что первый из сопоставляемых показателей выше второго. Для сравнений с оценками душили положительное Δ означает, что докладчик оценивает своё понимание выше, чем его оценивают со стороны. Объём выборки в сравнениях с участием «душили» меньше основного, поскольку не каждый докладчик получил внешнюю оценку в соответствующий момент. Корреляции показателей вмешательства с величиной ретроспективного пересмотра в таблицу не включены: ни одна из них не достигла уровня значимости (см. раздел «Вопросы душили»). Тренд сдвига по уверенности в наличии ответов проверен критерием Джонкхир – Терпстры для упорядоченной альтернативы; приведённое p подтверждено перестановочной проверкой (50 000 перестановок).

Ограничения

Наиболее существенное ограничение, что все три оценки понимания получены с помощью самооценки, следующим по значимости ограничение находится в плоскости дизайна эксперимента в котором отсутствует внутрисубъектные измерения доклад → мастер-класс. Ретроспективная оценка происходит в момент когда происходит оценка текущая, и может быть подвержена искажению. Роли во время игры чередовались, опыт спрашивания мог оказать влияние на собственную оценку независимо от полученных вопросов.

Основным результатом нашего исследования мы считаем воспроизводимость эффекта, обнаруженного Розенблитом и Кейлом, на принципиально ином материале. Участники эксперимента, объяснявшие свой замысел другому человеку, оценили свое изначальное понимание ниже, чем в момент формулирования в письменном виде перед началом взаимодействия, одновременно с этим, по мнению участников эксперимента, они стали понимать свой замысел лучше, чем понимали до разговора с внешним партнером.

На наш взгляд, именно наличие двух сдвигов – ретроспективного и сразу после взаимодействия – является содержательным результатом нашего исследования. Мы не можем утверждать, что участник стал понимать свои мысли лучше, а только то, что он изменил оценку, подняв уровень понимания после диалога с «душилой» и понизив уровень изначальной оценки в ходе ретроспективной оценки. Расхождение данных суждений составило 0,91 на CodeFest и 1,22 на PeopleSense и является самым ярким эффектом всего исследования ($d_z = 0,80$ и $0,93$).

Иначе говоря, разговор с внешним партнером сдвинул не столько понимание (дизайн эксперимента не предполагал измерения данного показателя), сколько то, как участник оценивает собственное понимание.

Проблематика данной работы рассматривает феномен пересмотра уровня глубины понимания субъектом исследования и фокусируется на вопросе, чем вызван пересмотр.

Изначальное, интуитивное предположение заключалось в том, что вопросами, то есть чем больше вопросов «душила» задал докладчику и чем вопросы существеннее, тем сильнее докладчик пересмотрит свою оригинальную оценку понимания, однако полученные данные этого предположения не подтверждают.

Первое. Сравнительное условие, при котором разработанный нами инструмент применялся без внешнего партнера, не произвело сдвига вообще: 5,11 против 5,19, $\Delta = +0,08$, $p = 0,380$, $d_z = 0,06$ при $n = 226$. При этом мы не можем сказать, что инструмент не отработал: в отношении собственного текста у 199 (88%) участников возникли вопросы, в среднем к четверти использованных ими в формулировках понятий. То есть вопросы к словам в тексте были, но оценка не сдвинулась, таким образом получается, что процедура сработала на первом этапе (спросить), но не сработала на втором (изменить оценку).

Второе. При спрашивании самого себя 36% участников ответили, что знают ответ практически на все вопросы, 19% знают ответы больше чем на половину и лишь 4% не знают ответов. Возможно, что при спрашивании самого себя, собственное сомнение снимается иллюзией собственного объяснения (понимания).

Третье. Даже если вопрос поступает извне, полученные данные указывают на то, что количество вопросов не предсказывает величину пересмотра изначальной оценки. В силу ограничения дизайна исследования мы не можем сказать, была ли зависимость между интенсивностью самой дискуссии и параметрами заданных вопросов. Ни количество вопросов, ни число «душил» на одного докладчика, ни величина субъективной полезности вопросов не связаны ни с ретроспективным пересмотром оценки, ни с приростом текущей оценки, все коэффициенты корреляции находятся в диапазоне от -0,21 до +0,28 и ни один не достигает уровня значимости.

Вместе три фактора указывают на один вывод: не инструмент имеет значение при пересмотре оценок, а внешняя позиция, источник, из которого вопросы поступают. Имеет значение перевод проверки из внутреннего плана во внешний. Оказывается, что такого перевода достаточно для изменения субъективной оценки понимания, дальнейшее наращивание интенсивности вопросов не добавляет ничего.

Прямое сравнение условий доклада и мастер-класса подтверждает данный вывод. Более того, оно указывает на то, что разговор ничего не изменяет в понимании, но помогает докладчику сдвинуть самооценки (ретроспективную и после сессии вопросов), сдвинуть их в разные стороны.

Полученный результат согласуется с критерием беглости суждения, о котором написано во введении к настоящему исследованию. Решение «мне понятно» в условиях доклада принимается не по результату содержания (ответов на вопросы), а по субъективной легкости, с которой мысль рождается в голове [Koriat, 1997]. Собственный текст обладает беглостью в максимальной степени: он только что написан, логика «понятна», и потому момент «прекратить спрашивать» срабатывает раньше, чем при вопросах извне. Проверка в собственной голове использует те же механизмы, которые и породили текст – механизмы, которые, по нашему мнению, связаны с воображением и иллюзиями.

Внешний собеседник «ломает механизм» не тем, что он знает больше, – как раз наоборот, скорее всего меньше, он обусловлен тем, что «душила» не имеет доступа к внутреннему контексту докладчика. Внешнему собеседнику требуется, чтобы автор проговорил вслух то, что он думает, что понимает. Перевод внутренней речи во внешнюю оказывается механизмом, при котором автор обнаруживает, что логика, изложенная им письменно, несостоятельна. Отсюда и природа невосприимчивости к количеству вопросов: докладчик обнаруживает «поломку»

логики на первом же ответе, дальше идет повторение уже сделанного открытия с новыми терминами и положениями, описанными текстом.

Этим мы объясняем отличие нашего исследования от исследования Розенблита и Кейла, в котором объектом объяснения выступало внешнее устройство: застёжка-молния, замок, спидометр, швейная машинка. Знания о внешних устройствах у респондентов исследования Розенблита и Кейла были фрагментарными и заимствованными. В нашем исследовании объектом выступает собственный замысел, к которому у респондента имеется эксклюзивное понимание и, что важнее, авторская позиция.

Эта разница, казалось бы, незначительная, на самом деле имеет в своей основе совсем иную природу: стоимость признания. Понять, что не понимаешь устройство бытового прибора, значительно «дешевле» психологически, чем признать тот факт, что не понимаешь, зачем приехал на конференцию за счет работодателя. Таким образом, можно сказать, что дизайн нашего эксперимента проверяет эффект в значительно более трудных условиях и в этих же условиях воспроизводит (на второй конференции).

Следует также отметить тот факт, что «душнилы» во всех замерах оценивают понимание докладчика ниже, чем он оценивает сам себя, и после обсуждения, казалось бы, разрыв должен был или исчезнуть, или уменьшиться, но этого не происходит. На CodeFest он составил $\Delta = 0,37$, $p = 0,027$.

Наше предположение, что обсуждение сблизит собеседников, не подтвердилось. Данные показали иное. Обе оценки в ходе обсуждения возрастают – и самооценка, и внешняя, – при этом дистанция между ними сохраняется. «Душила» оценивает докладчика выше, чем до обсуждения, но все равно ниже, чем докладчик оценивает себя.

Мы предполагаем, что обсуждение не сближает взаимные оценки, но при этом повышает взаимную удовлетворенность (активирует специфический социальный механизм): докладчик доволен, тем что ответил, а «душила», тем что получил ответ. Полезность вопросов оценена докладчиками в 5,83 и 5,81 балла, полнота ответов «душилами» – в 5,61 и 5,59. А также с очень высокой вероятностью часть разрыва несократима в принципе, потому что у докладчика остается доступ к несказанному, тогда как у «душнилы» его нет, и длительность обсуждения не способна компенсировать нехватку у «душнилы» «внутреннего знания» докладчика.

Следовательно, можно сделать вывод, что внешняя оценка не изменяет самооценку, не «повышает» понимание, а создает условия для появления более строгой точки отсчета самооценки понимания докладчиком.

Следует также отметить принципиальный факт: средние значения скрывают неоднородность результатов. Ретроспективную оценку понизили 40 из 89 участников CodeFest и 20 из 41 на PeopleSense, повысили 20 и 6, а не изменили 29 и 15 соответственно. Кроме того, у каждого шестого участника (17% на обеих площадках) обсуждение понизило текущую оценку понимания, что позволяет утверждать с определенной долей уверенности, что разговор с внешним спрашивающим («душилой») повышает вероятность увидеть в своих суждениях иллюзию понимания у половины участников, и некорректно делать утверждения, что разговор с внешним спрашивающим снижает иллюзию у всей экспериментальной выборки.

Вопрос, кто и по какой на самом деле причине принимает решение пересмотреть свою точку зрения, остается открытым и представляет наиболее перспективное направление дальнейшего исследования.

Обнаруженные в текущем исследовании эффекты нельзя называть большими: при пороге надежности обнаружения $d_z = 0,30$ и $0,45$ мы имеем $d_z = 0,32$ и $0,50$ для выборки $n = 88$ на CodeFest и $n = 41$ на PeopleSense, результаты лежат вблизи границы чувствительности

исследования, в связи с чем напрашиваются два вывода:

Воспроизведение эффекта на второй площадке, с вдвое меньшим числом участников, усиливает доверие к результату.

Отсутствие связей с показателями вмешательства нельзя трактовать как доказанное отсутствие: при $n = 88$ мы уверенно обнаружили бы корреляцию от 0,29 и выше, но не слабее.

Различие величин эффекта между площадками (0,43 и 0,68) мы отдельно не проверяли, а доверительные интервалы перекрываются почти целиком, поэтому содержательно интерпретировать различие мы не считаем возможным. Мы можем лишь отметить, что задания между площадками различались. На CodeFest участники обосновывали решение приехать на конференцию, а на PeopleSense – почему не применяют предложенную практику работы с вниманием. Литература, посвященная иллюзии понимания, указывает, что именно требование объяснить механизм, а не обосновать «почему / зачем?», дает эффект [Fernbach et al., 2013], что, на наш взгляд, является правдоподобной гипотезой, которую на текущих данных проверить не представляется возможным.

Тем не менее при максимально осторожной политике формулирования выводов одно следствие, на наш взгляд, обладает достаточной устойчивостью, чтобы быть сформулированным. Самопроверка понимания неэффективна не потому, что человек не использует вопросы, – наши данные указывают на то, что использует, – а потому что проверка использует тот же критерий/механизм, который и создает иллюзию.

Полученные результаты можно интерпретировать несколько скромнее, чем при первоначальном проектировании дизайна эксперимента: внешний источник вопросов и обязанность ответить вслух делают критерии самооценивания более жесткими. Ни количество вопросов, ни их качество не имеет значения, значение имеет наличие адресата. В контексте обучения каким-либо навыкам это означает, что короткий разбор в паре с кем-то несет в себе больше ценности, чем развернутый самостоятельный аудит.

Отметим, что внешняя оценка систематически ниже самооценки и остается ниже после разговора. Практический смысл подобного положения заключается в том, что мнение слушателя не следует воспринимать как окончательный вердикт о качестве написанного докладчиком.

Результат исследования позволяет сформулировать дальнейшее направление изысканий, и заключаться оно будет в поиске объективной меры измерения понимания, которая позволила бы отделить простой пересмотр ценности суждения от истинного изменения понимания. Дополнительно нам было бы интересно сопоставить (случайным образом) три условия: живой собеседник, самостоятельная оценка и автоматический спрашивающий в виде большой языковой модели. Данные настоящего исследования не позволяют прямую проверку сопоставляемости условий. Если существенна позиция, то языковая модель эффекта не даст, если существенен вопрос, то даст. И в завершение нам интересно узнать устойчивость эффекта, а именно: сохраняется ли результат пересмотра оценки спустя какое-то время или самооценка возвращается к исходной величине, как только внешний собеседник исчезает.

Заключение

В данном исследовании мы поставили вопрос воспроизведения иллюзии глубины понимания на материале, где у человека отсутствует заимствованное знание, – на его собственном замысле (мыслях, идеях, задачах, действиях), а также нам было интересно выяснить механизм, при котором происходит пересмотр самооценки, если вообще происходит.

Для поиска ответа мы использовали специально разработанную авторскую многопользовательскую веб-игру «Душнота по понятиям» и проверили практическую применимость инструмента в условиях двух отраслевых конференций.

С помощью многопользовательской игры были собраны данные посетителей конференции, решивших принять добровольное участие: 88 полных наборов самооценок на CodeFest и 41 на PeopleSense, а также 226 ответов в условиях, при которых участник задавал вопросы самому себе.

Эффект воспроизвелся на обеих площадках ($+0,43$ и $+0,68$, $p = 0,003$, $d_z = 0,32$ и $0,50$). После разговора с внешним спрашивающим участники оценили свое ретроспективное понимание ниже, чем оценивали в момент формулирования текста. Одновременно с этим текущая (постфактум) оценка выросла. Расхождение двух суждений (ретроспективного и постфактум) оказалось самым выраженным эффектом исследования ($0,91$ и $1,22$ по 7-ми бальной шкале при $d_z = 0,80$ и $0,93$). Однако следует отметить, что собранные данные позволяют зафиксировать именно пересмотр суждения о понимании, объективных мер измерения понимания дизайн эксперимента не предусматривает.

Тем не менее наиболее содержательным результатом оказался тот, которого мы не ожидали. Наше исходное предположение состояло в том, что величина пересмотра оценки определяется интенсивностью, то есть количеством и качеством заданных вопросов. Полученные данные это предположение не подтвердили. Ни число вопросов, ни субъективная полезность вопросов, ни число «душило» на одного докладчика не связаны с величиной пересмотра оценок. Все коэффициенты лежат в диапазоне от $-0,21$ до $+0,28$ и не достигают уровня значимости. При этом следует отметить, что при условии отсутствия собеседника (на докладе) сдвига не произошло вовсе, хотя инструмент выполнил свою функцию: в условиях доклада у 88% участников возникли вопросы к собственному тексту. Вопросы были – пересмотра оценок нет.

Из этого следует основной вывод работы – значение имеет не вопрос, а позиция, из которой он поступает. Перевод проверки из внутреннего плана во внешний, при котором имеется адресат, которому необходимо ответить вслух, оказывается достаточным условием пересмотра; любое дальнейшее наращивание интенсивности не добавляет ничего. А самопроверка оказывается неэффективной, потому что человек задает себе вопрос и оценивает ответ по критерию беглости, породившему исходную уверенность.

Мы также обнаружили, что разрыв между самооценкой и внешней оценкой не устраняется после обсуждения. Душила по-прежнему оценивает понимание ниже, чем докладчик. Мы полагаем, что часть этого разрыва несократима в принципе: у докладчика остается доступ к несказанному, у «душили» доступа нет и не будет.

Отдельно следует отметить, что эффект изменения оценки не однонаправленный и не является общим свойством выборки. Ретроспективную оценку понизили 45% и 49% участников, повысили 22% и 15%, не изменили 37% и 17%, а еще примерно у каждого шестого участника текущая оценка снизилась после обсуждения. Полученный результат описывает направление сдвига, но скрывает обстоятельства, при которых сдвиг оказался возможным. Вопрос о том, кто и по какой причине пересматривает свои суждения, остается открытым.

В практическом плане интерпретировать полученные в ходе эксперимента результаты можно следующим образом: короткий разбор в паре с внешним собеседником ценнее любого, даже самого развернутого, самостоятельного аудита собственных убеждений. При этом собеседнику не требуется ни обладать экспертностью в области обсуждаемой проблемы, ни задавать избыточное количество вопросов, требуется лишь позиция спрашивающего и

обязанность докладчика отвечать вслух. Вместе с тем внешнюю оценку не следует принимать в качестве какого-либо вердикта, она систематически ниже и остается ниже после разговора.

Дальнейшее направление работы мы видим в поиске объективной меры измерения понимания, которая позволила бы отделить изменение личностного суждения от изменения глубины понимания в прямом рандомизированном сопоставлении трех условий: при наличии собеседника, без собеседника и при ответах на вопросы большой языковой модели. А также в проверке устойчивости эффекта во времени, а именно: сохраняется ли пересмотренная оценка после исчезновения внешнего собеседника (человека или большой языковой модели).

Библиография

1. Дорофеев М. Путь джедая. Поиск собственной методики продуктивности. 5-е изд.
2. Агамалиев Р. Нам только кажется, что мы понимаем или объяснить что-то – это не борщ сварить // Цифровой Сад [Электронный ресурс]. URL: <https://rustamagamaliev.ru/garden/нам-только-кажется-что-мы-понимаем-или-объяснить-что-то-это-не-борщ-сварить/> (дата обращения: 04.08.2026).
3. Поварнин С. Искусство спора. Как читать книги. 1924.
4. Alter A.L., Oppenheimer D.M. Uniting the tribes of fluency to form a metacognitive nation // Personality and Social Psychology Review. 2009. Т. 13. № 3. С. 219-235.
5. Bartlett F.C. Remembering: A Study in Experimental and Social Psychology. Cambridge: Cambridge University Press, 2003. 344 с.
6. Farnsworth W. The Socratic Method: A Practitioner's Handbook. Boston: David R. Godine, Publisher, 2021. 264 с.
7. Fernbach P.M. [и др.]. Political extremism is supported by an illusion of understanding // Psychological Science. 2013. Т. 24. № 6. С. 939-946.
8. Hamilton K., Yao M. Blurring boundaries: effects of device features on metacognitive evaluations // Computers in Human Behavior. 2018. Т. 89.
9. Jordan K. [и др.]. Trivially informative semantic context inflates people's confidence they can perform a highly complex skill // Royal Society Open Science. 2022. Т. 9.
10. Koriat A. Monitoring one's own knowledge during study: a cue-utilization approach to judgments of learning // Journal of Experimental Psychology: General. 1997. Т. 126. № 4. С. 349-370.
11. Kruger J., Dunning D. Unskilled and unaware of it: how difficulties in recognizing one's own incompetence lead to inflated self-assessments // Journal of Personality and Social Psychology. 1999. Т. 77. № 6. С. 1121-1134.
12. Rozenblit L., Keil F. The misunderstood limits of folk science: an illusion of explanatory depth // Cognitive Science. 2002. Т. 26. № 5. С. 521-562.
13. Schwarz N. Metacognitive experiences in consumer judgment and decision making // Journal of Consumer Psychology. 2004. Т. 14. № 4. С. 332-348.
14. Sparrow B., Liu J., Wegner D.M. Google effects on memory: cognitive consequences of having information at our fingertips // Science. 2011. Т. 333. № 6043. С. 776-778.
15. Wegner D.M. Transactive memory: a contemporary analysis of the group mind / под ред. B. Mullen, G.R. Goethals. New York: Springer, 1987. С. 185-208.

The Illusion of Explanatory Depth on the Material of One's Own Design: Self-Diagnostics in a Group Using a Web Game

Rustam Agamaliev

Master's Degree Student,
National Research University "MIET",
124498, 1, Shokina square, Zelenograd, Moscow, Russian Federation;
e-mail: rustam.agamaliev@gmail.com

Maksim Dorofeev

Master's Degree Student,
National Research University "MIET",
124498, 1, Shokina square, Zelenograd, Moscow, Russian Federation;
e-mail: maxim.dorofeev@gmail.com

Abstract

The article examines the reproducibility of the illusion of explanatory depth effect (the illusion of understanding, when a person believes that they understand a subject more deeply than they actually do) on the material of the participant's own design, rather than an external technical device, as in the classic study by Rosenblit and Keil, where subjects overestimated their understanding of how household appliances work. The experimental intervention was implemented using the authors' multiplayer web game "Dushnota po ponyatiyam" ("Stuffiness by Concepts"), in which a participant recorded in writing their own statement according to a cause-and-effect template (formulating a certain concept or relationship), assessed the clarity of their design on a scale, answered "silly" questions from other participants (simple, unexpected questions testing the boundaries of understanding), and then assessed clarity again – current (how clear it is now) and retrospective (how clear it was at the beginning before the discussion), which allows recording a shift in self-assessment. Data were collected at two industry conferences, CodeFest and PeopleSense, with the participation of IT and design specialists: complete sets of self-assessments were obtained from 88 and 41 participants, and an additional 226 responses were collected in a comparative condition, where the same questions were applied by the participant to their own text without an external interlocutor, which made it possible to separate the effect of social interaction from the effect of self-reflection. Processing was performed using the paired Student's t-test with calculation of Cohen's d_z for estimating the effect size, the chi-square test for analyzing response distributions, and the Jonckheere–Terpstra test for identifying trends in ordinal data. It is shown that after the discussion, participants rate their previous understanding lower than at the time of writing the text (+0.43 and +0.68 points, $p = 0.003$), while the current assessment increases; the divergence of the two judgments (current assessment and retrospective assessment of previous understanding) turned out to be the most pronounced effect of the study, indicating that dialogue with other participants exposes the illusion of understanding. At the same time, neither the number of questions, nor their subjective usefulness, nor the number of questioners is related to the magnitude of the revision, and in the condition without an interlocutor, no shift occurs at all, which testifies to the key role of social interaction in triggering metacognitive revision. It is concluded that what matters is not the question itself (as a stimulus for reflection), but the external position from which it comes (the Other as a bearer of a different perspective, leading out of monological thinking), and the obligation to answer aloud, which forces one to articulate and test the boundaries of one's own understanding, making the game an effective tool for self-diagnostics and adult learning.

For citation

Agamaliyev R., Dorofeev M. (2026) Illyuziya ob'yasnitel'noy glubiny na materiale sobstvennogo zamysla: samodiagnostika v gruppe s pomoshch'yu veb-igry [The Illusion of Explanatory Depth on the Material of One's Own Design: Self-Diagnostics in a Group Using a Web Game]. *Psikhologiya. Istoriko-kriticheskie obzory i sovremennye issledovaniya* [Psychology. Historical-critical Reviews and Current Researches], 15 (6A), pp. 76-99. DOI: 10.34670/AR.2026.29.68.024

Keywords

Illusion of understanding, illusion of explanatory depth, self-assessment of understanding, retrospective assessment, metacognition, fluency of judgment, external interlocutor, web game, adult learning.

References

1. Dorofeev M. The Jedi Path. Searching for One's Own Productivity Method. 5th ed.
2. Agamaliev R. It only seems to us that we understand, or explaining something is not like cooking borscht // Digital Garden [Electronic resource]. URL: <https://rustamagaliev.ru/garden/нам-только-кажется-что-мы-понимаем-или-объяснить-что-то-это-не-борщ-сварить/> (accessed: 04.08.2026).
3. Povamin S. The Art of Argument. How to Read Books. 1924.
4. Alter A.L., Oppenheimer D.M. Uniting the tribes of fluency to form a metacognitive nation // Personality and Social Psychology Review. 2009. Vol. 13. No. 3. Pp. 219-235.
5. Bartlett F.C. Remembering: A Study in Experimental and Social Psychology. Cambridge: Cambridge University Press, 2003. 344 p.
6. Farnsworth W. The Socratic Method: A Practitioner's Handbook. Boston: David R. Godine, Publisher, 2021. 264 p.
7. Fernbach P.M. et al. Political extremism is supported by an illusion of understanding // Psychological Science. 2013. Vol. 24. No. 6. Pp. 939-946.
8. Hamilton K., Yao M. Blurring boundaries: effects of device features on metacognitive evaluations // Computers in Human Behavior. 2018. Vol. 89.
9. Jordan K. et al. Trivially informative semantic context inflates people's confidence they can perform a highly complex skill // Royal Society Open Science. 2022. Vol. 9.
10. Koriath A. Monitoring one's own knowledge during study: a cue-utilization approach to judgments of learning // Journal of Experimental Psychology: General. 1997. Vol. 126. No. 4. Pp. 349-370.
11. Kruger J., Dunning D. Unskilled and unaware of it: how difficulties in recognizing one's own incompetence lead to inflated self-assessments // Journal of Personality and Social Psychology. 1999. Vol. 77. No. 6. Pp. 1121-1134.
12. Rozenblit L., Keil F. The misunderstood limits of folk science: an illusion of explanatory depth // Cognitive Science. 2002. Vol. 26. No. 5. Pp. 521-562.
13. Schwarz N. Metacognitive experiences in consumer judgment and decision making // Journal of Consumer Psychology. 2004. Vol. 14. No. 4. Pp. 332-348.
14. Sparrow B., Liu J., Wegner D.M. Google effects on memory: cognitive consequences of having information at our fingertips // Science. 2011. Vol. 333. No. 6043. Pp. 776-778.
15. Wegner D.M. Transactive memory: a contemporary analysis of the group mind / ed. by B. Mullen, G.R. Goethals. New York: Springer, 1987. Pp. 185-208.