

АНДРЕЙ Г. КУЗНЕЦОВ

Европейский университет в Санкт-Петербурге; Университет ИТМО, Санкт-Петербург, Россия

ORCID: 0000-0002-0249-5890

НИКОЛАЙ И. РУДЕНКО

Европейский университет в Санкт-Петербурге, Россия

ORCID: 0000-0001-9511-3881

Как объяснить технологические аварии? «Человеческий фактор», социальный конструктивизм и онтологический поворот в изучении неопределенностей беспилотных автомобилей

doi: 10.22394/2074-0492-2021-4-119-146

119

Резюме:

В публичном дискурсе о технологических авариях доминирует популярное объяснение через «человеческий фактор». Оно делает существенное утверждение, что человек по определению склонен ошибаться, а машина — нет. В сфере беспилотного транспорта оно появилось в результате демодализации американскими СМИ, а затем стейкхолдерами БА результатов исследования Национальной администрации дорожной безопасности США 2008 года, которое утверждало, что водители — критическая причина 94% всех ДТП. В данной статье мы хотим показать, какие теоретические и социально-политические проблемы существуют с объяснением через «человеческий фактор». Для этого мы рассматриваем альтернативу в виде концепции технологической системы как конфликтующего набора правил, которые следуют за кон-

Кузнецов Андрей Геннадиевич — кандидат социологических наук, научный сотрудник, Центр исследований науки и технологий, Европейский университет в Санкт-Петербурге; аналитик, Центр научной коммуникации, Университет ИТМО. E-mail: akuznetsov@eu.spb.ru

Руденко Николай Иванович — кандидат социологических наук, научный сотрудник, Центр исследований науки и технологий, Европейский университет в Санкт-Петербурге. E-mail: diogenstyx@gmail.com

Исследование выполнено при поддержке гранта Российского научного фонда (проект РНФ № 20-78-10106) «Беспилотные автомобили и общество: взаимодействие технологий, социоэкономических сценариев и регулирования в радикальной инновации».

текстуализирующими практиками, которую предложил британский социолог Брайан Уинн. Мы сравниваем такое толкование с объяснением девиантного поведения, которое Роберт Мертон дал в работах 1930-х гг. Критикуя утилитаристов, Мертон показывает, что девиации вызываются противоречиями в социокультурной структуре общества. В обеих концепциях сбои предстают как результат реляционных эффектов напряжения и противоречия между элементами систем. В качестве иной, и более реалистичной альтернативы работе с авариями мы предлагаем идеи Аннемари Мол и Джона Ло. Последние, анализируя аварии, выделили четыре режима определения блага внутри споров после аварий: мобильная утопия, абсолютизм, менеджеризм и практические манипуляции. Мы показываем, что и объяснение через человеческий фактор, и теория Мертона, и отчасти STS склоняются к утопическим режимам (первым трем), в то время как последний режим, опирающийся на онтологический поворот, предлагает радикальный проект изменения режимов объяснения и обвинений в авариях: он позволяет артикулировать иные отношения между онтологиями аварий, сделать неутопистские версии технологий более реальными и публичными.

Ключевые слова: человеческий фактор, беспилотные автомобили, Брайан Уинн, Роберт Мертон, социальный конструктивизм, онтологический поворот, Аннемари Мол, Джон Ло

120

Andrey G. Kuznetsov¹

European University at St. Petersburg; Saint Petersburg State University of Information Technologies, Mechanics and Optics, Russia

Nikolay I. Rudenko²

European University at St. Petersburg, Russia

How to explain technological accidents? The “Human Factor”, Social Constructivism, and the Ontological Turn in Exploring the Uncertainties of Autonomous Vehicles

Abstract:

Public discourse about technological accidents is dominated by the popular explanation through the “human factor”. It makes the essential assertion that a human, by definition, is prone to error, but a machine is not. In the field of autonomous vehicles, it emerged as a result of first the US media—and subsequently, stakeholders—demodalizing the results of a 2008 US

1 Kuznetsov Andrey Gennadievich — PhD in Sociology, Research Fellow, Center for Science and Technology Research, European University at St. Petersburg; Tenured Associate Professor, Institute for International Development and Partnership, ITMO University, andrey.kuznetsov.29@gmail.com

2 Rudenko Nikolay Ivanovich — PhD in Sociology, Research Fellow, Center for Science and Technology Research, European University at St. Petersburg, diogenstyx@gmail.com

National Highway Safety Administration study that claimed drivers were the critical cause of 94% of all road traffic accidents. In this article, we want to show what theoretical and socio-political problems exist with an explanation through the “human factor”. To this end, we consider an alternative in the form of the concept of a technological system as a conflicting set of rules that follow the contextualizing practices proposed by the British sociologist Brian Wynne. We compare this interpretation with Robert Merton’s explanation of deviant behavior in the 1930s. Criticizing the utilitarians, Merton shows that deviations are caused by contradictions in the socio-cultural structure of society. In both conceptual schemes, failures are presented as the result of relational effects of tension and contradiction between the elements of the systems. For a different and more realistic alternative of dealing with accidents, we highlight the ideas of Annemarie Mol and John Law. The latter, analysing accidents, identified four modes of determining the good within disputes after accidents: mobile utopia, absolutism, managerialism, and practical manipulation. We show that both the explanations through the human factor, Merton’s theory of deviation—and, to some extent, STS—lean towards utopian regimes (the first three), while the latter regime, based on an ontological turn, proposes a radical project of changing the modes of explanation and accusations of accidents: this makes it possible to articulate different relationships between the ontologies of accidents, to make non-utopian versions of technologies more real and public.

121

Keywords: Human factor, autonomous vehicles, sociology of technology, Brian Wynne, Robert Merton, Bruno Latour

Введение

16 июня 2021 г. в Сан-Франциско “беспилотный” автомобиль (далее — БА) Waymo сбил человека, пересекавшего пешеходный переход на электросамокате. Человек получил повреждения мягких тканей тела, но покинул место происшествия самостоятельно. У БА слегка пострадала табличка с номером. Власти Калифорнии обязывают составлять протоколы всех происшествий с БА и делать их публичными. Из протокола этого столкновения [State of California 2021] узнаем подробности.

Автомобиль Waymo подъехал к перекрестку в “беспилотном” режиме¹, чтобы совершить поворот налево, и остановился на красный свет. В салоне в это время находился водитель-испытатель компании, который должен переключаться между двумя ролями. Когда машина движется в беспилотном режиме, такой водитель выполняет функцию супервайзора [Кузнецов 2022]. Эргономы и специалисты

1 В прямом смысле, т. е. беспилотным без кавычек, этот режим вождения назвать нельзя [см.: Кузнецов 2022].

по человеко-машинному взаимодействию говорят в этом случае, что человек следит за контуром управления (on the loop). Но если возникает нештатная ситуация или показания приборов аномальны, водитель может принять роль оператора, т. е. взять управление на себя, встроиться в контур управления (in the loop).

Когда загорелся зеленый, машина, оставаясь в *беспилотном режиме*, выехала на пешеходный переход, готовясь поворачивать налево, но начала пропускать встречный транспорт и других участников движения. Тем временем загорелся желтый свет, а БА Waymo все еще стоял, частично перекрывая перекресток и отказываясь поворачивать налево. В этот момент водитель-испытатель взял управление на себя и начал совершать поворот налево. Через четыре секунды автомобиль Waymo, находясь в *ручном режиме*, на скорости 6 миль в час (примерно 9,5 км/ч) сбил человека на электросамокате, пересекавшего на скорости 7 миль в час пешеходный переход на красный свет как для транспорта, так и для пешеходов.

Это событие служит хорошей отправной точкой для целей нашей статьи. Обошлось без серьезных последствий, и в последующем обсуждении превагируют не ажиотаж и эмоции, а трезвые попытки понять, что же произошло. Вряд ли за данной мелкой аварией последует публичное расследование. Но в одном из сообществ на онлайн-форуме Reddit, посвященном БА, возникла небольшая, но предметная и рациональная дискуссия [Reddit 2021], в ходе которой пользователи пытались интерпретировать протокол-описание инцидента. Обратимся к ее содержанию.

Был ли водитель нетерпелив в этой ситуации? Нет, утверждает один из пользователей форума, экс-водитель-испытатель Waymo, их обучают выводить автомобиль из беспилотного режима, когда того требует ситуация, как в данном случае. Был риск застрять на перекрестке и заблокировать его. Но из отчета ясно, что автомобиль не стоял посреди перекрестка, а лишь блокировал часть одного из переходов. Так что водитель мог подождать следующего зеленого сигнала светофора. Но экс-водитель-испытатель не сдается и предлагает не винить водителя в этом инциденте. Да, он принял неверное решение, но их работа связана с очень большим стрессом. Для одного из пользователей информация о том, чему обучают водителей-испытателей Waymo, звучит как инсайт. Она позволяет понять, почему во всех отчетах об инцидентах с Waymo, их автомобили находятся в ручном, а не беспилотном режиме.

Кто тогда виноват в столкновении — программное обеспечение Waymo или водитель? Пользователи дают несколько объяснений аварии через «человеческий фактор (далее — ЧФ)», в которых не алгоритмы или сенсоры, а именно человек виноват в инциденте. Согласно одной версии, водитель неправильно оценил ситуацию. Он

мог остаться на перекрестке на еще один цикл светофора, и ничего страшного не случилось бы. В другой версии выдвигается гипотеза, что сенсоры Waymo будто бы заметили электросамокат на тротуаре заранее, а алгоритмы предсказали его возможное поведение (проезд через переход на красный) и, вероятно, поэтому БА не собирался поворачивать, даже когда у него появилась такая возможность. Водитель не понял этой «прозорливости» БА, принял неверное решение и взял управление на себя. В еще одной версии ссылаются на неназванную книгу, написанную физиком, который разрабатывал системы безопасности для атомных реакторов. Основная мысль этой книги в том, что все аварии на атомных станциях случались в результате того, что люди нарушали протоколы безопасности, потому что ошибочно думали, что знают систему лучше.

Дискуссию в сообществе Reddit можно резюмировать так: проблема, скорее всего, в водителе. Автоматика сделала правильный выбор. Все это говорит о том, что компьютер очень мало или вообще не способствовал возникновению этого инцидента. Waymo просто нужно исправить некоторые баги, которые иногда приводят БА к путанице или делают его слишком нерешительным. Один из самых заплюсованных комментариев утверждает: «Это интересно, поскольку это тот случай, когда присутствие страховочного водителя делает его (БА) менее безопасным. Это может быть одним из ожидаемых знаков зрелости беспилотных автомобилей — водитель-испытатель становится самым большим источником рисков в беспилотном автомобиле» [Ibid.]. На периферии дискуссии оказываются замечания о том, что если у Waymo действительно есть видео, доказывающее, что их технологии все делали правильно, а проблема именно в водителе, то почему они их не выложат?

123

Что показывает этот пример?

1. Даже в такой, с виду простой, ситуации вовсе не очевидно, что именно произошло. Требуется интеллектуальная работа, чтобы понять и объяснить аварию. В данном случае, ее провели пользователи Reddit. Во-вторых, вышеописанная дискуссия показывает, как выглядит объяснение аварий через ЧФ на практике. Участники дискуссии несколькими разными способами пытались перенести вину за происшествие с алгоритмов Waymo на водителя-испытателя. Эту схему интерпретации мы назвали объяснением технологической аварии через ЧФ. Единственная попытка оправдать водителя исходит от пользователя, непосредственно знакомого с подобными ситуациями изнутри.

2. Несмотря на тривиальность (это всего лишь незначительная авария) и микроскопичность (это всего лишь маленькая дискуссия на Reddit), этот пример, обнажая механизмы объяснения через ЧФ, делает нас чувствительными к подобным схемам интерпретации

в лапидарных высказываниях других более заметных и статусных акторов.

Послушаем, например, руководителя Yandex Self-Driving Group Дмитрия Полищука: “[М]ашине нужно предсказать, что со всем этим [дорожной ситуацией] случится дальше. И вот здесь *человеческий фактор наиболее неопределенный* (курсив наш — АК, НР). Если бы все ездили строго по правилам, то предсказывать было бы довольно просто. Но типовая ситуация, когда кому-то нужно повернуть налево, он стоит в крайнем правом ряду и в последний момент начинает движение через всю разметку... И все это не подчиняется абсолютно никаким правилам, это просто происходит в ряде случаев” [Голыжина 2019].

К ЧФ апеллируют не только в России. Вот как обосновывает необходимость развития БА автогигант General Motors: “Представьте мир без автомобильных аварий. Наши БА нацелены искоренить *ошибку человеческого водителя* — главную причину 94% аварий, что приведет к меньшему количеству травм и смертей” (курсив наш — АК, НР) [General Motors 2021]. Исполнительный директор Waymo Текедра Мавакана приводит ту же самую статистику и добавляет: “[У] нас есть возможность инновировать и *убрать человека из уравнения*, что значит, что вы не полагаетесь на то, чтобы человек был на чеку в любой момент. Так что вам не нужна водительская лицензия, вы садитесь в машину, и водитель Waymo безопасно везет вас туда, куда вам надо” (курсив наш — АК, НР) [Recode 2021].

Но подобные повествования становятся драмой, когда речь идет об авариях с человеческими жертвами. Например, после аварии в марте 2018 г., в которой погиб Уолтер Хуан — владелец Tesla Model X, оснащенной функцией Autopilot, представители компании-производителя заявили, что “ложное впечатление, что Autopilot небезопасен, причинит вред другим [людям] на дороге... *эта авария могла произойти только если мистер Хуанг не обращал внимания на дорогу*” (курсив наш — АК, НР) [Baggu 2021].

Итак, мы видим, что в публичном дискурсе о технологических авариях и сбоях доминирует объяснение через ЧФ. Стейкхолдеры БА (СЕО, разработчики, полиси-мейкеры, аналитики) возлагают вину за аварии, ошибки, неэффективность в работе сложных социотехнических систем на человека (пользователя, потребителя, оператора и т. д.). У этого клише много вариаций: “человеческая ошибка”, “ошибка пользователя”, “непредвиденные последствия”, “оператор нарушил правила безопасности”, “пилот нажал не на ту кнопку, потому что забыл об особенностях этой модели самолета”, “инженер-испытатель смотрела в свой телефон, а не на дорогу”. В этих с виду простых заключениях человек определяется как «слабое звено» социотехнической системы, как «зона моральной деформации», веду-

щая к сбою системы [Elish, 2019: 41]. Эта нативная теория — важная часть обоснования и оправдания не только БА, но и других “умных” технологий.

Но отсылка к пресловутому ЧФ — не единственный способ объяснить технологические аварии. Чтобы начать думать о технологиях по-другому, нам не обязательно оправдывать людей и винить во всем машины. Нужно сосредоточиться на процессах взаимодействия между различными компонентами технологических систем, будь они людьми, механизмами или алгоритмами.

Цель этой статьи — показать, что объяснения через ЧФ проблематичны как в интеллектуальном, так и в социально-политическом плане. В качестве своего кейса мы выбрали публичные дискуссии об авариях и неопределенностях беспилотных автомобилей. Однако наш анализ применим и к другим технологиям. В работе мы последовательно решаем четыре задачи. Во-первых, мы рассмотрим технологию БА как социально организованную систему производства сюрпризов, одна из форм которых — аварии. Во-вторых, проследим истоки возникновения дискурса о ЧФ. В-третьих, артикулируем структурную гомологию аргументов социально-конструктивистской теории технологических аварий в исследованиях наук и технологий (далее — STS) и структурной социологией Р. Мертона. В-четвертых, отталкиваясь от этой гомологии, исследуем способ анализировать технологические аварии, учитывающий уроки онтологического поворота в STS.

125

Технология как социальный эксперимент: сюрпризы, неопределенности, аварии

Что такое технология? Мы, конечно, не намерены обозреть здесь всю традицию философии и эмпирических исследований технологий, чтобы дать на этот животрепещущий вопрос однозначный ответ на все случаи жизни и исследования. Мы лишь исследуем следствия одного частичного ответа, позволяющего понять, что сегодня происходит вокруг БА. *Технология — это социальный эксперимент* [Stilgoe 2018: 27].

Когда мы слышим историю об инциденте с Waymo, возникает вопрос: как этим летним вечером БА вдруг оказался на одном из перекрестков Сан-Франциско? Чтобы понять это, нам нужно кое-что узнать о разработке и тестировании БА вообще и о Waymo в частности.

Waymo нацелены на создание сервиса внутригородского беспилотного роботакси. Поэтому основной упор в нем делается на автоматизацию практик вождения в комплексной среде внутри американских городов. Важная часть этой автоматизации — ис-

пользование разных сред тестирования [Руденко 2022]. 2008 г. стал важной вехой в развитии современного автопрома. Стремясь оптимизировать расходы на тестирование (которые составляют до 30% от стоимости производства машины) на фоне финансового кризиса, автопромышленники стали широко использовать инструменты математической симуляции [Leonardi 2012]. Последние заменяют дорогостоящие тесты на дорогах и в лабораториях. Новые игроки типа Waymo, пришедшие в автопром из ИТ, естественным образом наследуют эту стратегию. Используя симуляторы типа XView и Carcraft, они виртуально проигрывают миллионы возможных сценариев на дороге, чтобы обучить ПО на основе искусственного интеллекта управлению БА [Madrigal 2017].

Роботакси Waymo должны освоить множество практик для езды в комплексной среде. Для этого подмножества практик вождения типичной проблемой является навигация пересечений на одном уровне или перекрестков разного типа. Поэтому Waymo тестирует свои БА на полномасштабных полигонах, моделирующих эту среду и многое из нее заимствующих. Например, компания построила на одном из своих полигонов точную копию перекрестков, часто встречающихся на дорогах общего пользования в США [Madrigal 2017].

Но таких лабораторных тестов все равно недостаточно. Реальные проезды по дорогам общего пользования — источник критически важных данных для лидеров разработки БА. Стратегическим ресурсом для компаний-разработчиков является получение разрешений от региональных и муниципальных властей на право тестировать свои незавершенные технологии в реальном пространстве и времени городской жизни. Получив в ряду других компаний разрешение от властей Калифорнии, Waymo, согласно вышеописанным приоритетам, сделала основной местностью своих публичных испытаний “комплексную городскую и пригородную среду” как выражение “наиболее неправилосообразного множества социальных феноменов за пределами машины” [Hind 2019: 11-12].

Итак, факт нахождения БА Waymo на перекрестке в Сан-Франциско и характер инцидента с его участием не случайны. Мы имеем здесь дело с тщательно подготовленным испытанием, экспериментом.

Если это эксперимент, то что испытывается? Разработчики «беспилотников» ожидаемо описывают эти заезды как технические эксперименты в общественном пространстве. По их словам, тестируются способности машин к автоматическому вождению в среде более непредсказуемой, чем закрытый полигон или симулятор. Но испытывается не только «техника», но и цифровые и материальные инфраструктуры, взаимодействия с другими участниками

дорожного движения, законы, нормативы, регуляторные режимы, общественное мнение и нравы. Эту тенденцию выводить на рынок не уже завершённые продукты, а как можно раньше начинать бета-тестирование в условиях, приближенных к реальным, превращая отдельные территории или сегменты коллективной жизни в “живые лаборатории”, называют новым режимом промышленной инновации [Laurent, Tironi 2015; Magres 2020]. В этом ко-креативном подходе к разработке инноваций тестируются не столько изолированные технологии, сколько общества целиком [Engels et al., 2019: 6]. Поэтому следует говорить именно о социальных экспериментах, а не о технических испытаниях в социальном контексте.

Если технология — это (социальный) эксперимент, то что такое эксперимент? В свое время историк науки Ханс-Йорг Райнбергер предложил понимать эксперименты как системы организованного производства сюрпризов [Rheinberger 1997]. В случае технологий такими сюрпризами могут быть аварии, нештатные ситуации, неопределенности или непредвиденные способы использования. В публичных экспериментах с БА техно-сюрпризы могут принимать разные формы: 1) курьезы вследствие неожиданного функционирования сенсоров или программного обеспечения БА, как то: алгоритмы Tesla распознали логотип “Бургер Кинг” как знак “стоп” [Rapiet 2020], а грузовик с отражающими сигналами как множество светофоров [LazyReplays 2021]; 2) мелкие происшествия типа разобранного нами инцидента с Waymo или практически идентичного случая с шаттлом Toyota, сбившего паралимпийца [MacInnes 2021]; 3) аварии, артикулирующие в одной драме судьбы людей, машин и компаний [см. Руденко 2018].

127

Здесь мы подходим к неожиданному следствию из нашего определения технологии. *Если технология — это социальный эксперимент, а эксперимент — это социально организованный способ производства сюрпризов в виде аварий, то технология — это социально организованная система производства аварий.* Разумеется, аварии — это лишь небольшая часть сюрпризов, производимых технологиями. Но далее мы сосредоточимся на них как самом ярком и неоднозначном типе техно-сюрпризов.

Аварии как проблема и ресурс для технологов и исследователей

Определение технологии как социально организованной системы производства аварий, конечно, очень ограниченно и имеет смысл только под определенным углом. Однако оно схватывает тривиальное, но редко осознаваемое положение дел. Всякий раз, когда мы говорим о технологиях, мы неявно подразумеваем аварии, поломки,

ошибки, сбои. То, что нет непогрешимых технических систем сколь неоспоримо, столь и тривиально.

Но формула “говорим технологии — подразумеваем аварии” указывает также на нормативную сторону техно-сюрпризов. Аварии есть только потому, что технологии — это нечто построенное, форма порядка, набор правил, предписывающих, позволяющих, запрещающих [Latour 1992b]. Только потому, что есть набор ожиданий того, что должно произойти, есть сюрпризы. Только потому, что есть набор правил, можно определить некоторые события или действия как ошибки. Только потому, что есть представление о штатной работе, можно определить нечто как аварию. Тогда наше определение технологии может стать менее вызывающим: *Технологии — это способы упорядочивания мира, которые специфически проводят границу между правильным и неправильным, и для этого порождают ситуации, в которых эта граница неожиданным образом испытывается на прочность.*

Получается, ни у разработчиков, ни у пользователей нет опции, при которой у них могут быть технологии, но не будет аварий. Но у них есть другие выборы. Какие техно-сюрпризы они желают произвести? И какие выводы они делают из них?

128

Отношение технологов к авариям амбивалентно. С одной стороны, сама идея испытания разработки подразумевает, что она подвергается риску. Зачем? Затем, чтобы в малом и контролируемом масштабе сгенерировать ошибки, которые можно затем исправить и сделать технологию более устойчивой и надежной. В этом отношении аварии — ресурс разработки. Технологии учатся на авариях. С другой стороны, когда техно-сюрпризы оказываются достоянием общественности, они нередко становятся проблемой. Громкие аварии могут не только нанести ущерб репутации и навлечь санкции со стороны регуляторов, но и вообще поставить крест на разработке.

Для исследователей технологий аварии также являются и ресурсом, и проблемой. С одной стороны, аварии — важный ресурс по двум причинам. Во-первых, аварии и поломки, ставшие достоянием широкой публики, выводят социальные эксперименты технологий из приватной сферы технологов на арену публичных дебатов. Во-вторых, аварии дают много ценной информации. Они показывают, какие правила и допущения формируют “черные ящики” технологий. Они также выявляют неопределенности, которые можно упустить из виду или просто проигнорировать, когда все работает хорошо. Поэтому аварии и следующие за ними официальные расследования и публичные разногласия — незаменимый источник информации и методологическая точка входа для исследователей.

С другой стороны, для исследователей, нацеленных на формирование рамок управления (нарождающимися) технологиями [Stilgoe et al., 2013], общественное восприятие аварий является проблемой.

Глубокий разрыв между практикой технологий и публичным дискурсом ведет к публичной ажитации и негодованию, а не рефлексии по поводу сложных причин поломок и катастроф. Это препятствует формированию социальной политики технологий, в основе которой лежит “социальное обучение” как способ институтов продуктивно осмыслять техно-сюрпризы [Stilgoe 2018: 26]. Разрыв между приватным и публичным заставляет технологов скрывать от публики практически любые неполадки вместо того, чтобы вовлекать ее в обсуждение ошибок. Это сокращает возможности для социального обучения [Wynne 1988: 158; Stilgoe 2018: 44]. Аварии как часть процесса обучения технологий остаются только приватным ресурсом отдельных компаний, а не условием для социального обучения и переопределения рамок управления технологиями.

Итак, технология — это социальный эксперимент, производящий неопределенности и сюрпризы, заметной формой которых являются аварии. Как разные акторы (технологи, регуляторы, полисмейкеры, (не)пользователи) обращаются с этими неопределенностями? Как они реагируют на аварии, объясняют их, распределяют ответственность за те или иные события? Эти вопросы открывают интересный путь для исследований технологий. Далее мы пройдем по нему немного. Объяснение аварий через ЧФ — один из способов обращаться с неопределенностями и сюрпризами технологий, который, как мы увидели выше, сегодня стал популярным клише публичного дискурса о БА и других технологиях. В следующем разделе мы рассмотрим происхождение этой схемы интерпретации. Сначала покажем, что в ее основе лежит намеренное упрощение сложности транспортной системы, мотивированное разработкой методологии анализа аварий с помощью понятия “критическая причина”. А затем — подчеркнем интеллектуальную и социально-политическую проблематичность объяснения через ЧФ.

129

Объяснение аварий через “человеческий фактор”: происхождение, популярность и проблемы одного клише в публичном дискурсе о технологиях

Стейкхолдеры БА приводят идентичные данные о доле аварий по вине ЧФ, т.к. ссылаются на исследование Национальной администрации дорожной безопасности США (National Highway Traffic Safety Administration, далее в тексте — NHTSA), проведенное в 2005–2007 гг. [National Motor Vehicle 2008]. Оно было посвящено выявлению основных факторов, ведущих к автомобильным авариям [Ibid.]. В США за громкой аварией обычно следует расследование, которое стремится объяснить событие и распределить ответственность

за него между участниками. В расследованиях технологических аварий подобие научного объяснения и судебной тяжбы становится очевидным. Они всегда порождают двойной вопрос: в чем причина и кто виноват? Аналитики NHTSA на основе многоуровневой выборки, учитывающей географию, день недели и время суток, создали из таких материалов репрезентативную для США базу данных аварий [Ibid.: 9].

В 2008 г. NHTSA приложила к исследовательскому отчету для Конгресса США статистическое резюме, в котором утверждалось, что “водители — критическое основание (reason) 2 046 000 аварий, что составляет 94% от общего числа” [Singh 2008: 1]. На каждую из остальных критических причин (транспортное средство, среда и неизвестные причины) приходится примерно по 2% аварий. Однако и в отчете, и в резюме исследователи NHTSA модализируют свои утверждения: “критическое событие и критическое основание (reason), которое привело к нему <...> не могут быть интерпретированы как причина (cause) аварии” [Ibid.: 2]. Эта модализация указывает на специфику их подхода. Для NHTSA авария — это линейная последовательность событий, анализируемая в терминах *критическое событие, предшествующее аварии и критическое основание* (reason) [National Motor Vehicle 2008: 10]. Первое обозначает крайнее событие в цепочке (поворот руля, отказ тормозов), которое привело к аварии. Второе — последнюю ошибку в причинной цепи, то, что непосредственно привело к критическому событию, предшествующему аварии. Цель NHTSA — выявить именно *критические основания* (reason), а не *причины* (cause) автомобильных аварий.

Однако крупные американские медиа отбросили ограничительные модальности и придали осторожным выводам NHTSA вид несомненных фактов [Zipper 2021]. Здесь начинается подъем объяснения через ЧФ в публичном дискурсе о БА. Стейкхолдеры БА мобилизуют стабилизированное и растражированное медиа утверждение “человек виновен в 94% всех автомобильных аварий” в своих нарративах, критикующих текущее состояние системы автомобильности и оправдывающих разработку и внедрение автоматизированного транспорта.

Но как возник сам подход, различающий причины и основания аварий? Проследим, как комплексное материальное событие аварии превратилось в набор дискретных элементов. NHTSA использовала подход Кеннета Перчонка из Корнеллской аэрокосмической лаборатории, который первым в 1970-х предложил четкую методологию исследования аварий [Perchonok 1972; Vinsel 2015: 873]. Он обрабатывал данные об авариях с помощью закрытого списка вопросов, на которые можно ответить только “да” или “нет” [Perchonok 1972: 4]. Это позволило применять дескриптивную статистику, использовать

четко очерченные понятия, направляющие внимание кодировщиков, и давать четкие основания для мотивированного заключения о том, почему произошла авария.

Как и NHTSA впоследствии, Перчонок отдавал себе отчет в том, что у аварии нет одной причины (cause). Авария — это множество взаимосвязанных причин и событий. Несмотря на это, он вводит ряд допущений, сильно упрощающих взгляд на систему автомобильности в целом и на ДТП в частности. Во-первых, он выдвигает постулат о том, что к авариям приводит ряд линейно связанных между собой событий. У любой аварии есть первое событие, промежуточные события и последнее событие. Последнее Перчонок называет “критическим событием”, так как оно “трансформирует ситуацию в такую, где авария неизбежна” [Ibid.: 7]. К примеру, водитель решил объехать идущий впереди автомобиль. Полагая, что ему хватит времени на маневр, он выезжает на встречную полосу и сталкивается с другим автомобилем. В этом случае маневр водителя считается критическим событием для аварии.

Во-вторых, Перчонок полагает, что у аварии может быть только одно критическое основание [Ibid.: 8]. Оно отвечает на вопрос: почему водитель сделал нечто, или автомобиль повел себя каким-то образом. К примеру, если маневр обгона сделан по причине того, что водитель не верно оценил пространство для маневра, то критическое основание заключается в ошибке принятия решения со стороны водителя [Ibid.: 15]. Почему основание может быть только одно? Ответ на вопрос дает понятие *виновной ответственности* (culpability) [Ibid.: 8]. До этого Перчонок рисовал сугубо механистический образ дорожного движения. Через понятие *виновной ответственности* он вводит свой анализ элементы социальной договоренности. Для него дорожное движение, как в виде ПДД, так и в виде неформальных правил на дороге, обязано своему нормальному функционированию коллективным ожиданиям водителей. Эти ожидания формируют правила поведения. Отклонение от правил есть нарушение ожиданий. Поэтому “водитель объявляется виновным, если его поведение создает аномальную ситуацию” [Ibid.]. Аномальная ситуация может возникнуть в ходе лишь одного действия — того, которое первым приводит к нарушению порядка. И таким событием является критическое событие.

Перчонок проводит пилотное исследование, которое демонстрирует продуктивность его подхода к анализу аварий. В частности, он утверждает, что 57% всех аварий связаны с человеком [Ibid.: 24]. К таким основаниям (reasons) относятся, например, неспособность человека принять правильное решение в определенных ситуациях или ограниченность его восприятия, не позволяющая принять во внимание все элементы дорожной ситуации. Еще около 30% всех

аварий случаются, когда человек совершает неправильные действия из-за влияния окружения. Например, когда машину “ведет” из-за скользкой дороги, или когда водитель не видит приближающейся машины из-за крутого поворота.

Как и NHTSA, Перчонк модализирует свои находки, указывает на их ограниченность. Но его подход оказался успешным настолько, что тридцать лет спустя им продолжают руководствоваться в NHTSA, потому что он заинтересовал стейкхолдеров системы автомобильности, создав в своем пилотном исследовании “пакет теории-метода” [Fujituga 1988], предлагавший аналитикам автоаварий сбор и анализ данных с помощью статистики и компьютера.

Важное следствие анализа Перчонка заключается в том, что он делает акцент на психофизиологических свойствах водителей. Такая трактовка схожа с характеристикой людей как нерадивых, ленивых, ограниченных, которая часто встречается в отчетах об авариях в сложных технических системах разного рода, а через них транслируется в публичный дискурс [Wynne 1988].

Что не так с объяснением через «человеческий фактор»?

132

Итак, клише ЧФ — одно из самых частых объяснений автомобильных аварий в публичном пространстве. Но что не так с этим объяснением? Во-первых, апелляции к ЧФ проблематичны в интеллектуальном плане. Они мешают понимать и исследовать технологии. За объяснениями через ЧФ может стоять наивный, но не безвредный посыл — “как было бы хорошо, если бы технологии полностью состояли из машин и исключали людей”. Такое объяснение техносюрпризов блокирует анализ технологии как сложной и внутренне противоречивой социотехнической системы. Оно неявно делает сущностное утверждение о природе человека и машины. Человек по определению склонен ошибаться, а машина — нет. Эта интеллектуальная установка препятствует пониманию того, что техносюрпризы и неопределенности вызваны не изолированными компонентами технологий, а взаимодействиями между ними, будь они людьми, механизмами или алгоритмами.

Во-вторых, объяснения через ЧФ заключают в себе проблему для политики технологий. С одной стороны, если акторы фокусируются на отдельных элементах социотехнических систем, сетей или ансамблей в качестве причин аварий, а не на анализе отношений между ними, то они не склонны реконфигурировать эти системы и регулирующие их институты. Апелляция к человеческим ошибкам как “зоне моральной деформации” препятствует изменению рамок политики инноваций и социальному обучению как способу

делать выводы из технологических аварий и неопределенностей на уровне институтов, а не отдельных людей и компаний [Stilgoe 2018: 44]. С другой стороны, объяснение через ЧФ может искусно использоваться стейкхолдерами БА и для критики текущего состояния системы автомобильности, и для оправдания неопределенностей и провалов своей собственной технологии. Лидеры индустрии БА мобилизовали для критики человеческого вождения утверждение аналитиков аварий, превращенное медиа в безусловный факт. Используя это упрощенное упрощение, они представляют процесс разработки БА как всего лишь замену человеческого водителя нечеловеческим. При этом неясно, где в этом упрощении третьего порядка проходит граница между стремлением играть по правилам публичной коммуникации и (само)обманом.

Но этот же аргумент от ЧФ используется для подспудного избавления стейкхолдеров БА от ответственности — “наши технологии не гарантированы от аварий, поскольку ненадежные люди неизбежно являются их частью”. Объяснения этого типа используются создателями БА и для оправдания аварий. Например, Yandex склонен квалифицировать происшествия со своими умными машинами как случаи, когда ими управляли инженеры-испытатели [Софрыгин 2019]. Это делается не только, чтобы вывести из-под критики саму технологию, но и чтобы минимизировать возможные судебные издержки в этой ситуации.

Таким образом, ЧФ оказывается стратегическим научно-судебным ресурсом компаний. Он уводит из-под критики саму технологию и возлагает весь груз ответственности на людей. Люди ошибаются всегда, машины — никогда. ЧФ оказывается идеальной критической уловкой — разработчики БА всегда правы [Latour 2004: 241].

133

Уинн и Мертон: гомология аргументов социального конструктивизма и структурной социологии

Хорошо, а как тогда объяснять технологические аварии, если не через ЧФ? Далее мы рассмотрим и соотнесем между собой аргументацию двух подходов: 1) социально-конструктивистский анализ технологий Б. Уинна, опирающийся на принципы социологии научного знания; 2) структурную социологию Р. Мертон. С одной стороны, это позволит нам начать думать об авариях в другом ключе, сосредоточившись на процессах взаимодействия между различными компонентами технологических систем, будь они людьми, механизмами или алгоритмами. С другой — структурное сближение предметно очень разных аргументов открывает перспективу новых артикуляций различий между STS и более традиционной социологией.

Социальный конструктивизм: аварии как следствие неправоподобности и внутренней противоречивости социотехнических систем

Еще в 1980–1990-х, до того, как клише ЧФ стало доминировать в критике и оправданиях БА, альтернативную концепцию предложили социально-конструктивистские STS. Поломки и аварии стали объяснять не ЧФ, а рассматривать как вызванные структурой или внутренней организацией социотехнической системы. Для этого надо было начать понимать технологии как внутренне сложные и противоречивые системы правил. Концепция Брайана Уинна хорошо репрезентирует эту традицию исследования. Уинн называет технические системы неправоподобными, потому что эксперты, которые их создают и поддерживают, не опираются на правила в своей деятельности, они создают правила, следуя за практиками. “Практики не следуют правилам, скорее правила следуют за развивающимися практиками” [Wynne 1988: 153].

134

Как технические системы становятся внутренне противоречивыми? Уинн приводит пример одной аварии. В мае 1984 г. недалеко от Ланкастера в подземном помещении гидротехнических затворов для перекачки воды произошел взрыв метана. Погибло 16 посетителей. В ходе расследования было установлено, что в тоннеле образовалась большая пустота, куда просочился и скопился метан из подземных угольных отложений. Затем при включении насосов газ заполнил помещения гидротехнических затворов под давлением воды и по неустановленной причине воспламенился. Почему это произошло? Отчасти потому, что операторы станции применяли неформальную практику, при которой промывные клапаны все время оставались открытыми. Эта практика противоречит официально установленной процедуре, которая требует держать клапан все время закрытым и раз в несколько недель проводить полную промывку, чтобы удалить скопившийся ил.

Почему операторы проявили такую “самодеятельность”? В действительности это было реакцией на протесты рыбаков, жаловавшихся на то, что после полной промывки их обычно чистая река мутилась в течении нескольких дней. Поэтому в качестве альтернативы официальной инструкции операторы начали медленно, но непрерывно смывать ил, оставляя клапан открытым. Почему никто не предупредил их о возможных последствиях такой эксплуатации системы? А потому что они создали для себя новое и неофициальное правило использования системы, о котором не знал никто, кроме них. Но даже если бы эксперты и знали об этом, то они явно не знали о метане и вряд ли могли бы предвидеть опасность образования пустоты в тоннеле.

Что произошло? Уинн называет это локальной “контекстуализацией” технической системы [Ibid.: 155]. В процессе эксплуатации даже изначально непротиворечивая система правил подвергается модификации и фрагментации из-за необходимости реагировать на новые ситуации, обстоятельства, локальные требования. Инженеры, операторы, пользователи не только следуют предустановленным правилам, но и генерируют новые. Как следствие техническая система становится внутренне противоречивой, что может (хотя и не обязательно) приводить к авариям.

Далее мы сопоставим критику ЧФ с критикой “человеческой природы” в социологической теории первой половины XX в. Это сопоставление позволит нам увидеть гомологию аргументов социально-конструктивистских STS конца 1980-х и теории социальной и культурной структуры конца 1930-х.

Роберт Мертон: объяснение моральных “аварий” через социальную и культурную структуру

В конце 1930-х гг. Роберт Мертон вслед за Толкоттом Парсонсом обрушился с критикой на “утилитаристскую” социальную теорию. В своей теории девиации он пытался объяснить, почему существуют поломки моральных, институциональных кодексов [Merton 1938]. Почему люди нарушают социально установленные правила? Мертона интересовали нарушения как неформальных и нравственных норм, так и формальных правовых законов. Весь спектр поведения, преподносящего социальные сюрпризы, он называл девиантным.

В объяснении девиации “утилитаристы” апеллируют к “человеческой природе”: у людей есть врожденные биологические влечения, которые прорывают ограничения, установленные социальным порядком. Задача общества — установить законы, которые будут либо эффективно сдерживать импульсы “человеческой природы”, либо управлять ими. Согласно “утилитаризму”, люди соблюдают закон либо, потому что это им выгодно, и тогда соблюдение законов — это результат рационального расчета, либо у них нет другого выбора, кроме как соблюдать правила, их поведение обусловлено социальным порядком, хотя бы они этого или нет [Ibid.: 672].

Обратим внимание на две вещи. Во-первых, “утилитаризм” проблематичен в социально-политическом плане. Всякий раз, когда мы сталкиваемся с преступлением как социальной поломкой, аварией, мы возлагаем ответственность за нее исключительно на индивида, нарушившего закон. Проблема всегда в отдельных людях и их “человеческой природе” и никогда в самом социальном порядке. Общество — всегда источник конформного поведения; “человеческая природа” — всегда источник девиации.

Во-вторых, аргумент “утилитаризма” подобен аргументу от ЧФ. Человек — источник аварии, потому что он по определению склонен ошибаться, а машина — нет. Социальная политика, основывающаяся на “утилитаристской” теории девиации, будет тяготеть к изоляции отдельных индивидов, а не к модификации социальной структуры. А политика технологий, которая руководствуется объяснением через ЧФ, будет стремиться “исключать человека из уравнения”, изолировать его как слабое звено системы. Стейкхолдеры БА, апеллирующие к ЧФ, по сути, говорят нам: “Изолируем нарушителей ПДД, поразим их в правах, и проблемы наших транспортных систем будут решены”. Другими словами, “утилитаристская” социальная политика будет блокировать обучение социального порядка точно так же, как это делает объяснение аварий через ЧФ в отношении технического порядка. Показательно, что в области законотворчества или правоприменения идеи, что для достижения большей конформности поведения членов общества следует просто изолировать преступников, или просто ужесточить наказание, давно дискредитированы. Сегодня они являются уделом лишь дилетантов. Но в публичном дискурсе о технологиях апелляция к ЧФ обладает большим хождением и авторитетом, она звучит из уст технологов и экспертов. Это позволяет нам представить, насколько мышление о технических порядках отстает от мышления социальных порядков.

136

Вернемся к Мертону. Что он предлагает в противовес “утилитаризму”? “Утилитаристы” дают асимметричное объяснение: конформное и девиантное поведение — следствие разных причин. Объяснение Мертона симметрично: источником *и* конформного, *и* девиантного поведения является социальная и культурная структура. Его объяснение основано на выделении в этой структуре четырех аналитически отличных, но эмпирически взаимодействующих элементов: *культурные цели*, *институциональные нормы*, *идеология*, *классовая структура* или социальная организация.

Культурные цели (КЦ) — это эмоционально окрашенный комплекс человеческих чаяний и мечтаний, которые наделены престижем, и потому желаемы индивидами. [Ibid.]. Примером КЦ может быть экономический успех. *Институциональные нормы* (ИН) “определяют, регулируют и контролируют приемлемые способы достижения [КЦ]” [Ibid.: 673]. Всякая культура, сцепляет между собой КЦ и ИН, необходимые для их достижения. ИН не обязательно совпадают с техническими нормами, которые с точки зрения индивида могут быть эмпирически более эффективными для достижения КЦ. Технические нормы определяют максимально широкий круг эмпирически возможных способов достижения КЦ. Но не далеко не все эти способы будут легитимны. Чтобы достичь КЦ культурно при-

емлемым образом, индивид должен ограничить свой выбор относительно узким кругом средств, определенных ИН.

Отношения КЦ и ИН варьируются вследствие операции “дифференциального акцентирования”. “Акцент на определенных [КЦ] может варьироваться независимо от степени акцента на [ИН]” [Ibid.]. На этом основании Мертон строит типологию социокультурных порядков — шкалу из трех положений, на полюсах которой находятся два противоположных, но несбалансированных типа, а между ними один сбалансированный тип социального порядка. Мертон делает акцент на типе порядка с диспропорциональным акцентом на некоторых КЦ без соответствующего акцента на ИН. Предельный случай такого порядка — ситуация, в которой граница между ИН и эмпирически эффективными нормами стирается, и выбор средств для достижения КЦ определяется “лишь техническими, а не институциональными соображениями” [Ibid.].

“Утилитаристы” считали, что цели человеческого поведения либо произвольны, либо являются следствием сугубо рационального расчета. Общество определяет лишь нормы достижения этих произвольных целей. В противовес этому Мертон утверждает, что не только нормы, но и цели, к которым должны стремиться люди, определяются культурной структурой. Более того, уже типология социокультурных порядков показывает, что несбалансированный акцент на одних ценностях в них нивелирует ценность других. Отсюда дисбаланс в культурной структуре может вести к моральным поломкам, девиации или нарушению правил этой же структуры.

137

Вводя две другие аналитические переменные — *идеологию* и *классовую структуру* — Мертон усиливает этот аргумент и показывает не только культурный, но и социальный генезис девиации и аномии. Например, *идеология* общества может требовать достижения экономического успеха от всех его членов, независимо от их положения в социальной структуре, а *классовая структура* — ограничивать доступ некоторых групп к ИН достижения этой КЦ. В результате индивиды в этих группах будут иметь дело с противоречащими друг другу правилами, и девиация среди них будет более вероятна, чем в других группах, не испытывающих напряжения между элементами социальной и культурной структуры. Таким образом, Мертон объясняет девиацию, моральные аварии, тем, что общество не является внутренне согласованной и непротиворечивой системой правил. Именно противоречия в социокультурной структуре общества ведут к девиации, а не прорывающие социальные ограничения импульсы “человеческой природы”.

Теперь несложно увидеть формальное подобие аргументов Мертона и социально-конструктивистских исследований технологий. Оба подхода отказываются от асимметричных и эссенциалистских

объяснений, усматривающих причины аварий (моральных или технических) в сущностных характеристиках отдельных элементов — “человеческая природа” или “человеческий фактор”. Оба подхода предлагают объяснения аварий как реляционных эффектов напряжения или противоречия между элементами социокультурной структуры или социотехнического ансамбля.

Онтологический поворот и множественные режимы учреждения блага в дискуссиях о технологических авариях

Вместе с тем, структурное подобие аргументов Мертона и Уинна позволяет обнаружить гомологию адресованных им в разное время критических замечаний, которые содержательно между собой никак не связаны. Так же как Деннис Ронг в 1960-х критикует Парсонса и Мертона за “сверхсоциализированную концепцию человека” [Wrong 1961], в начале 1990-х Бруно Латур выступает против столь же сверхсоциализированной концепции наук и технологий в социальном конструктивизме [Latour 1992a]. Ронг предлагает вернуться к более комплексной и диалектической концепции человеческой природы (в ее фрейдовском варианте), чтобы не девальвировать смысл гоббсовского вопроса о порядке и, как следствие, лучше понять сам механизм социализации. Латур предлагает разработать более реалистичную концепцию нечеловеческих акторов, нередуцируемых к социальным отношениям, значениям или инструментальности. С Латура в STS начинается онтологический поворот, который оформился как заметное движение к концу 1990-х — началу 2000-х. Далее мы рассмотрим, какие выводы из этой критики сделали ключевые представители этого движения — Аннемари Мол и Джон Ло.

Обратимся к их анализу материалов публичного расследования трех крупных железнодорожных аварий, произошедших в Великобритании с 1997 по 2000 гг. [Law, Mol 2002]. Работая в жанре эмпирической философии, Мол и Ло соединили этнографический интерес к практикам с философской заботой о благе. Вместо того, чтобы самостоятельно определять стандарты того, что такое хорошо и плохо, они исследовали как разные акторы (судьи, адвокаты, потерпевшие, менеджеры, железнодорожные работники) учреждали разные шкалы блага / худа в ходе судебного-научной дискуссии по поводу аварий. В результате они выделили четыре режима исполнения или учреждения (не)блага: три утопических и один — неутопический.

Утопизм — способ обращения с благом, который стремится освободиться от своей привязанности к конкретным локациям и к комплексным материальным практикам, и вместо этого размещает

себя в дискурсивном пространстве [Ibid.: 100]. Утопические режимы дискурсивны. Кроме того, они подразумевают возможность совершенства, ситуации, в которой нет напряжения или столкновения между различными благами [Ibid.: 84].

Первый из таких режимов они называют “мобильной утопией”. Это свободно плавающий дискурс, не связывающий себя необходимостью срочного действия. Для этого комплекса практик учреждения блага характерны обвинения, суть которых сводится к тому, что обвиняемые не соответствуют тому или иному высокому стандарту. Эта утопия мобильна, поскольку она не привязана ни к какому конкретному месту действия или практической задаче. Но она также не стремится дать и синопсис ситуации. Ее динамика представляет собой изолированные вспышки возмущения, в каждой из которых всего лишь демонстрируется, что благо, которое могло бы быть реализовано в данной ситуации, в действительности не было реализовано. Например, должны были соблюдаться инструкции безопасности, но они не соблюдались. При этом эта форма критики не принимает в расчет не только практические обстоятельства, но и другие типы благ и возможные напряжения между ними. Этнограф, собравший в одном месте и сопоставивший между собой эпизоды такого “праведного” и утопического гнева, может показать пустой характер этого режима учреждения блага.

139

Абсолютизм — второй утопический режим делания блага. Он также оперирует посредством обвинений. На этот раз в том, что какое-то определенное благо не было поставлено превыше всех других, хотя это следовало сделать. Например, безопасность и жизнь пассажира должна быть превыше прибыли, эффективности, устойчивости транспортной системы. Это утопизм, потому что не предполагается, что эта высшая форма блага может сосуществовать с другими благами. Исторически абсолютизм возник в ситуации, когда определенный авторитет мог навязать всем одно единственное благо [Ibid.: 90]. Сегодня, как показывают Мол и Ло, эту технику используют слабые акторы (в частности, пострадавшие в авариях и их представители), которые отчаянно стремятся быть услышанными среди множества голосов на публичных аренах. В этом смысле — это утопизм отчаяния.

Третий тип утопизма в обращении с благом, *менеджеризм*, существенно отличается от двух предыдущих. Он признает, что мир несовершенен, что в мире сосуществует много типов благ и они находятся в напряженных или противоречивых отношениях между собой. Он рассматривает нахождение баланса между ними как важную и сложную задачу. Однако это все равно разновидность утопизма. Сама процедура балансирования, управления противоречиями в виде аудита, оценки рисков, исследований и других форм испол-

нения подотчетности становится мета-благом. Менеджеризм — это скрытая утопия, где ценится рефлексивный и оправдательный дискурс [Ibid.: 94].

Однако задача Мол и Ло показать, что возможны и неутопические режимы делания (не)блага, которые не-дискурсивны, впутаны в специфичности практик, привязаны к конкретным локациям, имеют дело с настоящим (а не рефлексивной оценкой произошедшего), с материалами и телами (а не языком и мышлением). В своем исследовании они выделяют один такой режим задействия блага — *практические манипуляции* (tinkering) [Ibid.: 99]. Авторы показывают, что нормативная организация дискурсивного пространства публичного расследования аварии негативно сказывается на способности акторов, вовлеченных в этот режим, оправдать свою версию блага. Когда их призывают отчитаться в этом специфическом месте, их способ обращения с благом выглядит худом, ошибкой, некомпетентностью.

Этот режим артикулируется в материалах допросов на публичных слушаниях сигнальщиков, работавших во время одной из катастроф в электронном диспетчерском центре [Ibid.: 97]. Следствие пыталось выяснить, хорошо ли они выполнили свою работу? Могли ли они предотвратить или смягчить катастрофу? После того, как сигнальщики увидели, что поезд проехал на “красный свет” без авторизации, они должны были дать сигнал об экстренной остановке машинисту. По инструкции это должно было занять около двух секунд, но они подали сигнал где-то через двадцать с лишним секунд.

Работа сигнальщиков представляет собой непрерывный поток каждодневной рутины, вплетенной в конкретные локации и материалы: мониторы, документы, (ложные) сигналы, пути, шумы, смены, перерывы на обед и т. д. Эта работа предполагает молчание и срочное действие, а не говорение и рефлексия. Расследование пытается извлечь и перевести этот недискурсивный и телесный опыт в совершенно другую ситуацию. От сигнальщиков требуют отчета и оправданий: выпутаться из материальных практик своей работы и описать их простыми и ясными словами в утопическом дискурсивном пространстве. Их просят представить события в качестве линейной последовательности, которую затем следует дискурсивно разложить на несколько элементов, как то: установление проблемы, анализ ситуации, принятие решения и действие в соответствии с ним. Но у сигнальщиков нет языка для этого, они не могут дискурсивно дискретизировать свою практику так, чтобы это было приемлемо в публичном пространстве оправдания. Они не могут отчитаться перед следствием не столько потому, что стремятся уйти от ответственности, сколько потому, что их практика “развертывается неопределенно и относительно неподотчетно” [Ibid.: 99].

Итоговый посыл Мол и Ло заключается в том, чтобы найти способы публично говорить и репрезентировать режим практических манипуляций в его собственных терминах. Забота о качестве практик может помочь достичь блага там, где саморефлексия высокого модерна и требования подотчетности бесполезны [Ibid.: 101].

Вместо заключения

Какие выводы мы можем сделать из онтологического поворота в STS и, в частности, из исследования Мол и Ло? Во-первых, можно сопоставить рассмотренные нами объяснения поломок (моральных или технических) с режимами делания блага, выделенными Мол и Ло. Объяснения через ЧФ и утилитаризм тяготеют к утопическому абсолютизму. Они грезят совершенными порядками, организованными вокруг единой шкалы благ, и видят в склонности человека ошибаться или в дикой человеческой природе источник абсолютного зла. Структурная теория Мертона артикулирует различия между абсолютизмом и менеджериализмом. Есть сбалансированные и разбалансированные социальные порядки, и есть конформисты, которые справились (managed) с противоречиями благ в социальной и культурной структуре, и девианты, которым это не удалось. Социальный конструктивизм Уинна артикулирует различия между менеджериализмом и практическими манипуляциями. С одной стороны, практики оценки технологий и рисков живут идеей, что можно разработать научный метод принятия решений и создания сбалансированной системы правил безопасности, по которым будут функционировать технологические системы. С другой, STS показывают, что действительные практики разработки, поддержания и управления технологиями подразумевают контекстуализацию и нормализацию, они скорее порождают правила, чем следуют им. Именно это ведет к фрагментации технологических систем и создает потенциал для аварий.

141

Во-вторых, из этого следует, что объяснения технологических аварий через ЧФ, социальное конструирование и социокультурную структуру необязательно рассматривать как взаимоисключающие перспективы. Они сосуществуют. Дело не в том, что перспектив на аварии много, а в том, что это больше чем перспективы. Это способы делания событий авариями и распределения ответственности за них между причастными акторами. У каждого объяснения есть своя онтология. И если у STS есть задача альтернативно объяснять аварии, то значит их задача также состоит в создании новых онтологий технологических поломок. Как сделать эти объяснения более реальными, чем имеющиеся? Как могут или должны быть артикулированы различия между ними? Это предмет дальнейших исследований.

В-третьих, характеристики режима практических манипуляций — ситуативность, телесность, материальная комплексность, недискурсивность, непрозрачность и относительная неподотчетность — это неотъемлемые черты практик и человеческого, и автоматизированного вождения [Киселева 2022]. Предложенный Мол и Ло анализ предательства логики режима практических манипуляций при переводе его в дискурсивный формат публичного расследования позволяет лучше понять историю происхождения клише ЧФ. Оно возникло из понятия “критического основания”, которое появилось в результате судебно-научной практики дискретизации непрерывного потока событий автоаварии. Перчонку потребовалась серия допущений и упрощений, чтобы сделать неправилосообразные практики подотчетными в публичном пространстве. Элементы среды, маневры водителей должны были быть извлечены из практической комплексности дорожной ситуации, чтобы путешествовать в офисы NHTSA и конгрессменов. Но аргументация и выводы аналитиков автоаварий должны были дополнительно быть демодализированы и утопизированы медиа, чтобы стать в итоге клише ЧФ в публичном дискурсе о медиа. Но эти же процессы наблюдаются не только в публичных оправданиях БА, но и в дискуссиях о непрозрачности нейронных сетей — самом сердце этой технологии. Требования раскрыть “черные ящики” нестрогих алгоритмов и сделать их подотчетными, игнорируют непрозрачность и неправилосообразность технологий. А вся дискуссия тяготеет к оппозиции между метафорами просвещения и обскурантизма [Кузнецов 2020].

142

Все это говорит о необходимости изменить описанный режим объяснения, обвинения и подотчетности [Латур 2012] в публичных дискурсах и деланиях БА и других технологий. Для этого нужно будет не только предложить новые перспективы и объяснения, но и артикулировать иные отношения между онтологиями аварий, сделать неутопистские версии технологий более реальными и публичными.

Библиография / References

Голыжбина О. (2019) Дмитрий Полищук, «Яндекс»: «Абсолютно гарантировано, что беспилотное такси будет дешевле обычного». *Реальное время*. <https://realnoevremya.ru/articles/148125-pochemu-bespilotniki-yandeksa-do-sih-por-ne-vyshli-na-dorogi>.

— Golyzhbina O. (2019) Dmitry Polishchuk, Yandex: “It is absolutely guaranteed that an autonomous taxi will be cheaper than a regular one.” *Realnoe vremya*. <https://realnoevremya.ru/articles/148125-pochemu-bespilotniki-yandeksa-do-sih-por-ne-vyshli-na-dorogi>. — in Russ.

Киселева М. (2022) Аффигирования, артикуляции, обучение: к техноантропологической симметризации тел водителей и беспилотных автомобилей. *Этнографическое обозрение*, 1 (в печати).

— Kiseleva M. (2022) Affections, articulations, training: towards the techno-anthropological symmetrization of the bodies of drivers and autonomous vehicles. *Etnograficheskoe obozrenie*, 1 (in press). — in Russ.

Кузнецов А. (2020) Туманности нейросетей: “черные ящики” технологий и наглядные уроки непрозрачности алгоритмов. *Социология власти*, 32 (2): 157-182.

— Kuznetsov A. (2020) Nebulae of neural networks: “black boxes” of technologies and demonstrative lessons of the opacity of algorithms. *Sociology of Power*, 32 (2): 157-182. — in Russ.

Кузнецов А. (2022) Симметричная антропология технологий: от культур к коллективам мобильностей в описании беспилотных автомобилей. *Этнографическое обозрение*, 1 (в печати).

— Kuznetsov A. (2022) Symmetrical anthropology of technologies: from cultures to collectives of mobility in the description of autonomous vehicles. *Etnograficheskoe obozrenie*, 1 (in press). — in Russ.

Латур Б. (2012) Политика объяснения: альтернатива. *Социология власти*, 8: 113-143.

— Latour B. (2012) The Politics of Explanation: An Alternative. *Sociology of Power*, 8: 113-143. — in Russ.

Латур Б., Вулгар С. (2012) Лабораторная жизнь. Конструирование научных фактов. Глава 2. Антрополог посещает лабораторию. *Социология власти*, 1 (6-7): 178-234.

— Latour B., Woolgar S. (2012) Laboratory Life. The Construction of Scientific Facts. Chapter 2. An Anthropologist Visits the Laboratory. *Sociology of Power*, 1 (6-7): 178-234. — in Russ.

Руденко Н. (2018) Четыре смерти Элен Херцберг. *Археология русской смерти*, 6: 97-99.

— Rudenko N. (2018) The four deaths of Elaine Herzberg. *Archeology of Russian Death*, 6: 97-99.

Руденко Н. (2022) Симметричная антропология тестирования беспилотных автомобилей. *Этнографическое обозрение*, 1 (в печати).

— Rudenko N. (2022) The Symmetrical Anthropology of Testing Autonomous vehicles. *Etnograficheskoe obozrenie*, 1 (in press). — in Russ.

Софрыгин А. (2019) Первое ДТП Беспилотника Яндекс на общих дорогах. Виноват водитель или беспилот? *Bespilot.com*. <https://bespilot.com/news/508-yandex-crash>

— Sofrygin A. (2019) The first accident of a Yandex autonomous vehicle on public roads. The driver or the vehicle to blame? *Bespilot.com*. <https://bespilot.com/news/508-yandex-crash>. — in Russ.

Barry K. (2021) Elon Musk, Self-Driving, and the Dangers of Wishful Thinking: How Tesla's Marketing Hype Got Ahead of Its Technology. *Consumer Reports*. <https://www.consumerreports.org/automotive-industry/elon-musk-tesla-self-driving-and-dangers-of-wishful-thinking-a8114459525/>

Elish M. (2019) Moral Crumple Zones: Cautionary Tales in Human-Robot Interaction. *Engaging Science, Technology, and Society*, 5: 46–60.

Engels F., Wentland A., Pfotenhauer S. M. (2019) Testing future societies? Developing a framework for test beds and living labs as instruments of innovation governance. *Research Policy*, 48 (9): 1–11.

Fujimura J. H. (1988) The molecular biological bandwagon in cancer research: Where social worlds meet. *Social Problems*, 35 (3): 261–283.

General Motors (2021) *Safety*. <https://www.gm.com/stories/self-driving-cars>.

Halsey III A. (2019) More Americans have died in car crashes since 2000 than in both World Wars. *The Washington Post*. https://www.washingtonpost.com/local/trafficandcommuting/more-people-died-in-car-crashes-this-century-than-in-both-world-wars/2019/07/21/0ecc0006-3f54-11e9-9361-301ffb5bd5e6_story.html

Hind S. (2019) Digital navigation and the driving-machine: supervision, calculation, optimization, and recognition. *Mobilities*, 14 (4): 401–417.

Latour B. (1992a) One More Turn after the Social Turn: Easing Science Studies into the Non-Modern World. McMullin E. (ed.) *The Social Dimensions of Science*. Notre Dame, Ind.: University of Notre Dame Press: 272–292.

Latour B. (1992b) Where are the missing masses? The sociology of a few mundane artifacts. W. Bijker, J. Law (eds) *Shaping technology/building society: Studies in sociotechnical change*. Cambridge, Massachusetts; London, England: the MIT Press: 225–258.

Latour B. (2004) Why Has Critique Run out of Steam? From Matters of Fact to Matters of Concern. *Critical Inquiry*, 30 (2): 225–248.

Latour B., Fabbri P. (2000) The Rhetoric of Science: Authority and Duty in an Article from Exact Science. *Technostyle*, 16 (1): 115–134.

Laurent B., Tironi M. (2015) A field test and its displacements. Accounting for an experimental mode of industrial innovation. *CoDesign*, 11 (3–4): 208–221.

Law J., Mol A. (2002) Local Entanglements or Utopian Moves: An Inquiry into Train Accidents. *The Sociological Review*, 50 (1_suppl): 82–105.

LazyReplays (2021) Tesla runs into several traffic lights while over 120km/h. *YouTube*. https://www.youtube.com/watch?v=cO_tOq4W68w

Leonardi P. M. (2012) *Car Crashes Without Cars: Lessons about Simulation Technology and Organizational Change from Automotive Design*. Cambridge, Mass.: MIT Press.

MacInnes P. (2021) Toyota pauses Paralympics self-driving buses after one hits a visually impaired athlete. *The Guardian*. <https://www.theguardian.com/technology/2021/aug/28/toyota-pauses-paralympics-self-driving-buses-after-one-hits-visually-impaired-athlete>

Madrigal A. (2017) Inside Waymo's Secret World for Training Self-Driving Cars. *The Atlantic*. <https://www.theatlantic.com/technology/archive/2017/08/inside-waymos-secret-testing-and-simulation-facilities/537648/>

Marres N. (2020) What If Nothing Happens? Street Trials of Intelligent Cars as Experiments in Participation. S. Maasen, S. Dickel, C. Schneider (eds.)

TechnoScienceSociety: Technological Reconfigurations of Science and Society. Nijmegen: Springer/Kluwer: 111-130.

Merton R. K. (1938) Social Structure and Anomie. *American Sociological Review*, 3 (5): 672-682.

National Highway Transport Safety Agency (2021) *The Evolution of Automated Safety Technologies*. <https://www.nhtsa.gov/technology-innovation/automated-vehicles-safety>

National Motor Vehicle Crash Causation Survey (2008). *Report to Congress*. U.S. Department of Transportation. National Highway Traffic Safety Administration. <https://crashstats.nhtsa.dot.gov/Api/Public/ViewPublication/811059>.

Perchonok K. (1972) *Accident cause analysis. Final Report*. Cornell Aeronautical Laboratory, Inc. Prepared for U.S. Department of Transportation, National Highway Traffic Safety Administration.

Rapier G. (2020) Tesla's Autopilot confused a Burger King sign for a stop sign. The fast-food chain turned it into an ad. *Business Insider*. <https://www.businessinsider.com/tesla-autopilot-mistakes-burger-king-stop-sign-new-ad-2020-6>

Recode (2021) Waymo Co-CEO Tekedra Mawakana. *YouTube*. <https://www.youtube.com/watch?v=xaYCT7woLMo>.

Reddit (2021) Waymo DMV Collision Report on June 16th Collision with a Scooter (Waymo was in manual mode for 4 seconds before the crash). *Reddit*. https://www.reddit.com/r/SelfDrivingCars/comments/ofs9hc/waymo_dmv_collision_report_on_june_16th_collision.

Rheinberger H. J. (1997) *Toward a History of Epistemic Things: Synthesizing Proteins in the Test Tube*. Stanford, California: Stanford University Press.

Singh S. (2008) Critical Reasons for Crashes Investigated in the National Motor Vehicle Crash Causation Survey. *Traffic Safety Facts. Crash Stats*. Department of Transportation. National Highway Traffic Safety Administration. <https://crashstats.nhtsa.dot.gov/Api/Public/ViewPublication/812115>.

Stilgoe J. (2018) Machine learning, social learning and the governance of self-driving cars. *Social Studies of Science*, 48 (1): 25-56.

Stilgoe J., Owen R., Macnaghten P. (2013) Developing a framework for responsible innovation. *Research Policy*, 42 (9): 1568-1580

State of California. Department of Motor Vehicles (2021) *Report of traffic collision involving an autonomous vehicle*. https://www.dmv.ca.gov/portal/file/waymo_061621-pdf.

Vinsel L. (2015) Designing to the test: performance standards and technological change in the US automobile after 1966. *Technology and culture*, 56 (4): 868-894.

Wrong D. H. (1961) The Oversocialized Conception of Man in Modern Sociology. *American Sociological Review*, 26 (2): 183-193.

Wynne B. (1988) Unruly Technology: Practical Rules, Impractical Discourses and Public Understanding. *Social Studies of Science*, 18 (1): 147-167.

Zipper D. (2021) The Deadly Myth That Human Error Causes Most Car Crashes. *The Atlantic*. <https://www.theatlantic.com/ideas/archive/2021/11/deadly-myth-human-error-causes-most-car-crashes/620808>.

Рекомендация для цитирования:

Кузнецов А. Г., Руденко Н. И. (2021) Как объяснить технологические аварии? “Человеческий фактор”, социальный конструктивизм и онтологический поворот в изучении неопределенностей беспилотных автомобилей. *Социология власти*, 33 (4): 97-146.

For citations:

Kuznetsov A. G., Rudenko N. I. (2021) How to explain technological accidents? The “Human Factor”, Social Constructivism and the Ontological Turn in Exploring the Uncertainties of Autonomous Vehicles. *Sociology of Power*, 33 (4): 97-146.

Поступила в редакцию: 23.10.2021; принята в печать: 19.11.2021

Received: 23.10.2021; Accepted for publication: 19.11.2021