

Семинар "Педагогика ИИ"

Дата и время: 09.07.2025, 10.00

Место проведения: НИУ «МЭИ», г. Москва, ул. Красноказарменная, д. 13с3, аудитория М-101.

Участники: Роман Сергеевич Куликов (НИУ «МЭИ»), Олег Вадимович Глухов (НИУ «МЭИ»), Павел Романович Варшавский (НИУ «МЭИ»), Константин Александрович Петров (НИЦ "Курчатовский институт" - НИИСИ), Людмила Георгиевна Голубкова (НИУ «МЭИ»), Александр Павлович Еремеев (НИУ «МЭИ»), Артем Анатольевич Птицин (Институт Инноваций и Права), Павел Юрьевич Анучин (НИУ «МЭИ»), Сергей Викторович Вишняков (НИУ «МЭИ»), Маргарита Андреевна Орлова (НИУ «МЭИ»)

Повестка:

1) Тема основной дискуссии «Сравнение экспертного и нейросетевого подходов для решения прикладных и отраслевых задач»:

- Доклад «Лучшие практики применения инструментов ИИ для решения прикладных и отраслевых задач».

2) Выбор тематики следующего семинара.

Олег Вадимович Глухов:

Коллеги, рад всех приветствовать на очередном семинаре по педагогике искусственного интеллекта.

Напомню наш формат: это открытый семинар, в котором участвуют представители заинтересованных организаций. Семинар проводится на регулярной основе. Ведётся аудиозапись, после чего готовится стенограмма.

Предлагаю по традиции представиться. Особенно это касается новых участников. Меня зовут Олег Вадимович Глухов. Я ассистент кафедры Радиотехнических систем МЭИ, и являюсь организатором этого семинара.

Прошу всех участников представиться — пойдём против часовой стрелки.

Константин Александрович Петров:

Меня зовут Петров Константин Александрович. Я представляю НИИ системных исследований (НИИСИ).

Очень надеюсь стать проводником к нужным ресурсам для реализации больших проектов по ИИ: сначала их найти, а затем использовать для решения актуальных задач.

Людмила Георгиевна Голубкова:

Добрый день. Меня зовут Людмила Георгиевна Голубкова. Я представляю МЭИ. В рамках данного семинара я выступаю как потребитель его результатов. Надеюсь, что итоги нашей работы окажутся значимыми и полезными, позволят ИРЭ (Институт радиотехники и электроники) МЭИ и другим дружественным институтам и организациям обрести новые качества, в том числе и в области педагогических практик — как традиционных, так и основанных на искусственном интеллекте.

Павел Романович Варшавский:

Коллеги, меня зовут Павел Романович Варшавский. Я заведующий кафедрой прикладной математики, а с 2019 года — также и кафедрой искусственного интеллекта (МЭИ).

Можно сказать, что я являюсь носителем компетенций в данной области, поскольку и кафедра, и я лично занимаюсь этой проблематикой достаточно давно и глубоко.

На протяжении многих лет у нас проходил семинар по проблемам искусственного интеллекта, где поднимались научно-исследовательские и фундаментальные вопросы. Однако тематика данного семинара отличается — она затрагивает вопросы педагогики ИИ, и здесь уже инициатива в большей степени принадлежит Александру Павловичу Еремееву, который курирует данное направление в рамках Российской ассоциации искусственного интеллекта. Я постараюсь в своём докладе также кратко коснуться этих вопросов.

Александр Павлович Еремеев:

Меня зовут Александр Павлович Еремеев. Я профессор в области математики и искусственного интеллекта. Занимаюсь вопросами, связанными с ИИ, и считаю эту тему крайне важной, поскольку участвовал в её развитии с самого начала.

К сожалению, возможности искусственного интеллекта зачастую переоцениваются. Недавно один из специалистов в области систем безопасности ИИ справедливо отметил, что ожидания от подобных технологий во многом завышены — создаётся впечатление, будто ИИ способен заменить человека и решать любые задачи. Это не соответствует действительности.

Поэтому я считаю, что наш семинар чрезвычайно полезен. Подобные обсуждения уже проводились в рамках Научного совета Российской ассоциации искусственного интеллекта. Там также подчёркивалось: искусственный интеллект — это хороший помощник, но он не заменит человека, особенно в серьёзных и творческих задачах.

Роман Сергеевич Куликов:

Я Роман Сергеевич Куликов, директор Института радиотехники и электроники МЭИ. На этом семинаре выступаю с позиции заказчика. Мне необходим практический результат нашей работы, который можно будет внедрить в учебный процесс. В частности, задача — формирование искусственных ассистентов для инженеров с целью повышения производительности труда и автоматизации тех операций, которые поддаются автоматизации в рамках инженерной деятельности.

В краткосрочной перспективе планирую обучение студентов проектированию радиоэлектронных систем и их элементов с грамотным и продуктивным применением инструментов, основанных на искусственном интеллекте. Чтобы это стало возможным, необходимо сначала научиться обучать самих преподавателей и выстроить методологию. В этом контексте я рассматриваю себя как потребителя результатов данного семинара.

Артем Анатольевич Птицин:

Меня зовут Артём Анатольевич Птицин, я ведущий разработчик Института инноваций и права. Мы занимаемся внедрением искусственного интеллекта и его применением в тех сферах, где это действительно целесообразно. Придерживаемся принципа рациональности: стараемся не использовать ИИ ради самого ИИ, а внедряем его там, где он приносит практическую пользу.

Павел Юрьевич Анучин:

Меня зовут Павел Юрьевич Анучин, я инженер кафедры Радиотехнических систем (МЭИ). В настоящее время в своей работе я применяю методы ИИ и хотел бы больше погрузиться в данную область.

Олег Вадимович Глухов:

Коллеги, сегодня у нас содержательный доклад от Павла Романовича, предлагаю сегодня сосредоточиться преимущественно на нем — обсуждении лучших практик применения искусственного интеллекта для прикладных и отраслевых задач. У коллег накопился значительный практический опыт, которым они готовы поделиться.

И небольшая ремарка: при рассмотрении кейсов хотелось бы, чтобы каждый из них сопровождался ответами на вопросы — как ставилась задача? откуда она возникла? каким образом решалась и к какому результату в итоге удалось прийти?

Павел Романович Варшавский:

Коллеги, хотел бы поднять один вопрос. Понимаю, что на первый, установочный семинар мы не попали — в это время как раз проходила защита магистров из Китая. Наверное, я готов начать с доклада, но перед этим хотелось бы сказать несколько слов по теме, особенно по вопросу педагогики. Это направление нас также напрямую касается. Мы уже задумываемся о внедрении соответствующих решений. Руководство проявляет заинтересованность и, условно говоря, уже готово «купить» систему полностью. Однако, конечно, в этом процессе имеются и подводные камни.

Как уже отмечал Александр, многие вопросы обсуждались и на других семинарах. Здесь тоже наблюдаются определённые проблемы, и, думаю, будет полезно затронуть их в заключительной части моего доклада — для последующего обсуждения.

Доклад сегодня — наш совместный, в том числе с Александром Павловичем Еремеевым. Поскольку он не делегировал мне выступление полностью, думаю, в

процессе обсуждения он также подключится с комментариями и ответами на вопросы. Тем более что мы сегодня, по сути, представляем разработки научной группы под руководством Александра Павловича. Я также могу отнести себя к этой группе.

Начать хотел бы с краткой исторической справки. Кафедра прикладной математики занимается вопросами, связанными с искусственным интеллектом, уже достаточно давно — ещё с тех времён, когда, как говорится, проходили «огонь, воду и медные трубы». Причём на какой стадии сложнее — на начальной или на текущей — это ещё вопрос. Но в любом случае кафедра прошла все основные этапы становления в этой тематике.

Следует отметить, что истоки этих работ связаны с именем Дмитрия Александровича Поспелова — одного из пионеров в области ИИ и преподавателя нашей кафедры. Его научную школу в настоящее время продолжает профессор Еремеев.

Все мы — выходцы из научной школы, тесно связанной с системами поддержки принятия решений. Именно в этом контексте искусственный интеллект выступает как помощник, а не как замена человека. Хотя, конечно, есть разные точки зрения на то, что считать искусственным интеллектом — вопрос остаётся открытым и неоднозначным.

Например, наши коллеги из Российской ассоциации ИИ указывали на нюансы перевода термина *artificial intelligence*. В дословном переводе речь скорее идёт о «разумном рассуждении», но закрепилось выражение «искусственный интеллект». Это порождает неоднозначные трактовки, особенно в массовом восприятии. Однако мы, как правило, понимаем под ИИ научное направление в области компьютерных наук, направленное на моделирование разумных, интеллектуальных функций человека средствами программирования и математики.

Сегодня на волне популярности — нейробиомиметические подходы, которые действительно дают серьёзные результаты. В то же время в истории были периоды, когда акцент делался на системные методы, на разработку экспертных и продукционных систем. Кто-то делит ИИ на «сильный» и «слабый», но без «слабых методов», по сути, и о «сильных» говорить невозможно.

Кафедра прошла все этапы: от становления в 90-е, когда под руководством Дмитрия Александровича Поспелова началась подготовка специалистов в области интеллектуальных систем, до сегодняшнего дня. Это была одна из первых кафедр в СССР, а затем и в РФ, которая занималась именно этой проблематикой. Мы продолжаем развивать заложенные теоретические модели и методы.

Если переходить к примерам, то начнём с разработок научной группы под руководством Александра Павловича Еремеева. Речь идёт об инструментальном комплексе СИМПР (Инструментальная система моделирования принятия решений), в основе которого лежит аппарат таблиц решений. Это программная реализация систем поддержки принятия решений, позволяющая моделировать предметные области и ставить задачи на их основе. Язык таблиц решений оказался удобным для формализации непроцедурных знаний, что делает его понятным и доступным даже для пользователей без навыков программирования.

Таблицы решений позволяют удобно описывать процесс принятия решений, в том числе в сочетании с симуляторами и встроенными системами. По сути, это продукционные правила: «если – то». Таблица содержит условия и соответствующие действия, которые активируются при выполнении условий. Есть также расширенные возможности — например, проверка полноты базы знаний, обеспечение непротиворечивости.

СИМПР как инструментальное средство позволяет моделировать процесс принятия решений в рамках экспертных систем, где знания предметных экспертов переносятся в набор правил и могут использоваться без их непосредственного участия.

Сегодня основное внимание уделяется нейросетям, но экспертные системы по-прежнему актуальны — особенно там, где требуется объяснимость. Один из примеров использования таблиц решений — диагностика и контроль состояния турбогенераторов электростанций. В частности, выявление трещин в роторах и обрывов болтов. Система формализует параметры, которые необходимо контролировать, и предоставляет обоснованные рекомендации оператору.

Важно, что такие системы можно применять в разных режимах работы объекта. Они позволяют проводить анализ проблемных ситуаций и принимать решения о ремонте или управляющих воздействиях заранее. Вместо того чтобы вручную поднимать регламент или обращаться к эксперту, можно положиться на формализованное знание, встроенное в систему.

Правила в системе — продукционного типа. Они могут учитывать степень уверенности в результате. Например, в таблице явно прописаны отслеживаемые параметры, целевые показатели, и на этой основе выносится решение. В одном из применений было реализовано 11 правил, каждое из которых соответствует определённому диагнозу или управляющему воздействию.

Например, система может выявить трещины в роторе в режиме выбега. А в режиме работы под нагрузкой — предикторы их возможного появления. Это важные обобщающие выводы, которые строятся на основе анализа всей базы знаний.

Преимущество таблиц решений — в их прозрачности. Можно построить дерево решений, которое наглядно показывает, на основании каких условий система пришла к тому или иному выводу. Это недоступно нейросетям, в которых отсутствует объяснительный компонент. Именно поэтому экспертные системы ценны: у них всегда есть модуль объяснения, без которого невозможно применять ИИ в критических системах — будь то медицина или промышленность.

Если система рекомендует лечение или действие, всегда возникает вопрос: почему? И экспертная система способна на него ответить.

В дополнение — СИМПР также автоматически проверяет базу правил на непротиворечивость и полноту. Это позволяет заранее исключить логические ошибки и только после этого интегрировать систему. Пример использования — учебный процесс, где студенты записывают таблицы решений, а затем связывают их с симуляторами, например, на базе Simulink, моделируя, скажем, проезд перекрёстка.

Результаты тестирования, отладки и оптимизации также предусмотрены — система гибко масштабируется и адаптируется под задачу.

Роман Сергеевич Куликов:

Сейчас на слуху в основном нейросетевые подходы — о них говорят много, и действительно, результаты порой впечатляют, особенно когда они производят вау-эффект. Однако возникает вопрос: если сравнить эти нейросетевые подходы с экспертными системами, с теми же таблицами решений, то каковы их плюсы и минусы? Какие у каждого подхода сильные и слабые стороны?

Павел Романович Варшавский:

Ну, в принципе, один из важных плюсов я уже обозначил. Здесь мы получаем полностью интерпретируемые и обосновываемые решения. В таких системах не возникает вопроса неоднозначности. То есть, если система была грамотно спроектирована, и мы корректно смоделировали объект или процесс, то проблема объяснительного компонента просто не стоит — он встроен в саму архитектуру.

Более того, подобные системы — это именно то, что сегодня стараются интегрировать в понятие «доверенного ИИ». То есть не просто получить результат, как это бывает с большой языковой моделью, когда она что-то генерирует, и мы не можем понять почему, а дополнительно сопоставить её вывод с базой формализованных правил. Это даёт возможность провести верификацию вывода, проверить, насколько он согласуется с экспертными знаниями, и при необходимости — объяснить пользователю, почему именно такое решение предложено. То есть интеграция этих подходов тоже возможна.

Роман Сергеевич Куликов:

То есть, если говорить о плюсах и минусах — то плюс экспертных систем с таблицами решений заключается в стопроцентной прослеживаемости. Мы всегда можем точно определить, почему было принято то или иное решение, и на основании каких правил оно возникло.

В случае нейросетевых или более сложных моделей, наоборот, процесс принятия решения не поддаётся полной трассировке. Результат возникает в виде перемножения матриц и работы обученного софта — по сути, он формируется автоматически на основании обучающей выборки. Это усложняет интерпретацию и снижает прозрачность.

В чем тогда минус экспертных систем?

Павел Романович Варшавский:

Минус следующий, что если вы обратите внимание, то на каком-то этапе человек нужен для нейросетей — для формирования этого большого объёма данных под обучение. Хотя сегодня мы просто накидываем кучу материалов, которые система, самообучаясь, уже генерирует в рамках большой модели.

Здесь же вы видите (у экспертных систем), что при расстановке этих правил нужен эксперт-человек. То есть на уровне инженера по знаниям нужен кто-то, кто это заложит.

Роман Сергеевич Куликов:

То есть человек-эксперт нужен и в том, и в другом случае?

Павел Романович Варшавский:

Я бы так сказал: если сегодня смотреть на уже, скажем так, большие языковые модели — многое уже дано "на движке". То есть, мы только, например, если нам нужно GigaChat дообучать энергетическим вопросам, мы обращаемся к экспертам, чтобы они дополнили эту базу знаний. Но, с другой стороны, если база знаний уже накоплена, то дальше активное привлечение эксперта вам, по большому счёту, для нейросети не требуется.

Роман Сергеевич Куликов:

Это понятно, в нейросетях есть фаза эксплуатации и в экспертных системах...

Павел Романович Варшавский:

Нет, имеется в виду именно фаза разработки. А вот в фазе эксплуатации хочется отметить, что если предположить, что у нас ситуация на объекте будет как-то кардинально меняться, и появятся какие-то параметры, которые мы, например, до этого не учли...

Нейросетевая модель будет пытаться найти их решение. А здесь (в экспертных системах) вы пойдёте к человеку и скажете: «Вот, мы не учли, например, ещё какой-то показатель в роторе, да?» — Это означает, что придётся модель всё-таки дорабатывать и переделывать: и базу правил, и параметры, и многое другое.

Роман Сергеевич Куликов:

Что лучше, когда ИИ говорит «у меня непредсказуемое поведение» или пытается что-то сделать?

Павел Романович Варшавский:

Главное ограничение экспертных систем — это ситуация, когда в предметной области возникают случаи, по которым специалисты не могут дать однозначного заключения. Это может быть связано либо с высокой сложностью технической системы, либо с тем, что при анализе последствий неочевидно, как именно сработали отдельные механизмы. В таких случаях даже в критической ситуации эксперт не всегда готов сформулировать четкое и уверенное решение — зачастую появляется множество конкурирующих гипотез.

В этом плане нейросетевые модели обладают определённым преимуществом: они способны работать даже при неполноте данных. Их задача — сгенерировать возможные варианты ответов или гипотезы на основе входных параметров. Это, по

сути, инструмент без эмоциональной компоненты — если данные есть на входе, система даст результат на выходе.

Аналогично, когда эксперты не знают, как поступить, они организуют мозговой штурм. Нейросеть в такой ситуации можно использовать как вспомогательный инструмент генерации гипотез.

Роман Сергеевич Куликов:

Если сравнивать два подхода при управлении сложным объектом — один с использованием экспертной системы, другой — с нейросетевой моделью...

Павел Романович Варшавский

Качество экспертной модели будет всегда выше. Ключевое различие заключается в следующем: в случае экспертной системы каждое правило формализовано и имеет конкретного автора. То есть за каждым решением стоит определённый специалист, и в случае ошибки можно точно установить его источник. В нейросетевой системе такое отследить невозможно: результат формируется на основе внутренней структуры модели, и установить, кто конкретно «ответственен» за ошибку, невозможно.

Роман Сергеевич Куликов:

Хорошо, фиксируем. Второй плюс экспертной модели – идентифицируемая ответственность, который вытекает из прослеживаемости обработки данных (т.е. высокой уровень объяснимости). А какой-то минус у экспертной модели есть?

Константин Александрович Петров:

Позвольте, я выскажусь. У экспертных моделей очевидные ограничения. Один из основных минусов — это жесткость структуры. Если в базе знаний содержится, например, 10 или 11 правил, и позже возникает необходимость учесть двенадцатое, то приходится пересматривать всю систему: дополнять таблицу решений, привлекать эксперта, актуализировать логику.

При этом, чтобы определить, чего именно не хватает, сначала нужно выдвинуть гипотезы. И здесь как раз уместна роль нейросетевых моделей: они способны предложить возможные варианты, которые эксперт не смог бы сформулировать самостоятельно. На основе таких предварительных гипотез эксперт может доработать базу знаний, введя недостающее правило.

Таким образом, экспертная система обладает высокой точностью и контролируемостью, но ограничена в гибкости и требует ручного сопровождения при расширении.

Сергей Викторович Вишняков:

Коллеги, если обобщать, то экспертные системы оптимальны в ситуациях, когда предметная область может быть формализована в виде набора правил. Например, образовательный процесс: поведение студента регламентируется нормативными актами. На основе этих положений можно задать формальные правила, которые

автоматически определяют его переход на следующий курс, право на стипендию, продление обучения и т.п. Такая система способна заменить административного куратора, потому что структура процессов и условий формализована и известна...

Роман Сергеевич Куликов:

Но если происходит изменение законодательства — например, Госдума принимает новый закон...

Сергей Викторович Вишняков:

Тогда мы пересматриваем правила, обновляем базу, и система продолжает функционировать в соответствии с новыми условиями.

В отличие от этого, нейросетевые подходы целесообразно применять в случаях, когда система плохо формализуема и явных правил не существует. Мы лишь располагаем эмпирическим опытом, наблюдениями и историческими данными. Тогда обучение модели производится на выборке примеров, и правила извлекаются в скрытом виде в процессе обучения.

Павел Романович Варшавский:

Коллеги, все верно вы говорите, но поймите. Экспертные системы бывают двух типов:

1. Статические, где база знаний фиксирована;
2. Динамические, способные переключаться между различными наборами правил в зависимости от контекста или режима работы системы.

В текущем примере рассматривалась именно статическая модель. Возможности динамического переключения между базами знаний в данной реализации не демонстрировались.

Рассмотрим ситуацию, когда объект функционирует в нескольких режимах. В штатном режиме система опирается на один набор правил, тогда как при переходе в аварийное состояние — автоматически активируется альтернативная база знаний. Такой подход соответствует существующим техническим регламентам на многих объектах: существуют предписания как для штатных, так и для нештатных сценариев. Динамическое переключение между наборами правил реализуемо и целесообразно — это обеспечивает адаптивность системы к различным условиям функционирования.

В некоторых кейсах, о которых далее пойдет речь, классический экспертный подход оказывается ограниченным (т.е., где база знаний фиксирована). Несмотря на то, что изначально задача формулировалась в рамках классической экспертной парадигмы, оказалось, что не все аспекты можно эффективно описать формализованными правилами. Это обстоятельство стало основанием для внедрения дополнительных компонентов — об этом я расскажу на следующем примере.

Если готовы, предлагаю перейти к следующему кейсу.

Александр Павлович Еремеев:

Когда мы говорим об использовании систем или методов искусственного интеллекта, на первый план выходят два ключевых фактора:

1. Скорость генерации решения,
2. Пояснимость результата.

Мы совместно с коллегами разрабатывали прототип системы принятия решения реального времени. Это была интеллектуальная поддержка оператора при управлении подсистемами атомного энергоблока в штатном режиме. Также помощь осуществлялась в условиях перехода объекта из штатного режима в нештатный, до активации защитной автоматики (которая, напомним, обычно просто снижает нагрузку или отключает объект).

Для этих целей особенно удобен язык таблиц с решением — он ориентирован на матричные операции. Мы, например, использовали его ещё в кооперации с немецкими коллегами из ТУ «Ильменау» — при программировании станков с ЧПУ. Модификация решений — оперативная и простая. Студенты моделировали, например, сценарии для беспилотников на перекрёстке: стоит изменить один параметр — и модель сразу перестаёт сталкиваться. Правила можно быстро обновить для перехвата нестандартных ситуаций. Причём система заранее не запрограммирована на эти сценарии.

Есть два варианта:

- использовать накопленные знания из класса прецедентов — если такой прецедент уже был, можно применить готовое решение;
- если прецедент ближайший, то эксперт его адаптирует или введёт новое действие.

Но это уже не в реальном времени.

Здесь важно понимать ещё один момент: 90% базы знаний составляют нормативные документы. Мы, например, сотрудничали с НИИКА и использовали порядка 15 000 производственных правил, основанных на официальной документации. Только около 10% (а может, и меньше) добавлены экспертами.

Кстати, у атомщиков считается, что даже нештатные ситуации системе должны быть учтены. Однако, как показала Фукусима, предусмотреть абсолютно всё невозможно. Накладываются маловероятные события, что создаёт риски. Поэтому система должна быть способна действовать эффективно при таких маловероятных отклонениях.

Что ещё критически важно — это пояснимость решений. Когда у нас есть дерево решений — всё становится понятно даже непрофессионалу.

Как правильно отметил наш Президент, среди приоритетов ИИ — генеративные модели и большие языковые модели. Однако с этим связано явление галлюцинаций. По данным, например компании OpenAI, на некоторых задачах ошибки в ответах достигали 40–60%. Причина — нейросеть не объясняет своё решение.

Максимум — она показывает входной и выходной векторы. Но для критических ситуаций этого недостаточно. Мы должны понимать, почему система предлагает именно это решение.

Это и привело к постановке задачи разработки доверенного ИИ, где нейросети способны объяснять свои действия. Сейчас в них начали встраивать элементы производственных правил, что повышает прозрачность, но ухудшает быстроедействие и увеличивает время обучения. Тем не менее, пока эта интеграция работает не так хорошо. Это своего рода компромисс: если мы говорим о системах, где ошибка может

стоит очень дорого, следует использовать максимально прозрачные архитектуры, где мы видим процесс принятия решения.

Известны случаи, когда нейросеть принимала ошибочное решение, что приводило к судебным разбирательствам на Западе. Была ситуация, когда на face control человек тысячу раз распознавался корректно, а на тысяча первый — не распознавался. Например, человека снимали с рейса и уводили куда надо, и потом выяснилось, что всё было нормально – система дала сбой, но из-за этого возник стресс и прочие последствия. Такие случаи — пример скрытых рисков нейросетевых технологий. Да, они полезны, но и с ними нужно быть осторожными, это важно учитывать.

Роман Сергеевич Куликов:

Одной из целей семинара — выработать обоснованный подход к применению всего арсенала искусственного интеллекта, каким бы он ни был, в наиболее рациональной форме — для создания интеллектуальных ассистентов инженерам. Эти ассистенты должны реально повышать производительность и эффективность их труда. Вы так защищаете эти экспертные модели, как будто я на них нападаю. Я просто не специалист в них, поэтому хотел бы узнать. Попробуем, с другой стороны зайти, чтобы по пятому кругу не доказывать из чего состоит экспертная модель – я с первого раза понял.

Допустим, перед нами стоит задача — обучить тополога микросхем. Сегодня квалифицированных разработчиков микросхем остро не хватает: кто-то уехал за границу, кто-то переориентировался, а работы — в избытке. В России наблюдается кадровый голод при высоком спросе. Есть гипотеза: если учебная программа изложена в учебниках и включает дисциплины, базирующиеся на физике, программировании и других областях, а ежегодно по стране такие программы проходят несколько сотен студентов, то можно использовать инструменты искусственного интеллекта, чтобы создать интеллектуального ассистента разработчику.

Этот ассистент получал бы от разработчика задачу, предлагал решения, а разработчик, в свою очередь, контролировал бы результат. Это применимо к большинству инженерных профессий, где не требуется оперативная реакция в критической ситуации, как у хирургов или пилотов. Мы говорим именно о технических специальностях, не затрагивая пока творческие профессии — актёров, художников, архитекторов.

Гипотеза такова – создать интеллектуального ассистента, если использовать формализацию, обучающую выборку и другие подходы. Возникает вопрос: какое место в такой системе займут экспертные системы, а какое — нейросетевые методы, в частности языковые модели (LLM)? Как это, с вашей точки зрения, может быть скомпоновано?

Павел Романович Варшавский:

На мой взгляд, ответ достаточно очевиден. Я уже об этом говорил: интеграция методов ИИ даёт преимущества. Отказ от нейросетей или LLM нецелесообразен,

особенно с учётом того, как хорошо они работают с синтезом естественного языка. Это мощный инструмент, особенно в роли помощника. Но экспертный компонент, особенно в виде базы знаний и правил, должен формулироваться не нейросетью, а специалистом-предметником.

Например, эксперт может использовать LLM, чтобы ускорить подготовку черновиков правил, но всю ответственность за достоверность всё равно несёт человек. Я не могу сказать, что у нас есть такие примеры, по типу «Дайте на два учебника, мы быстро оцифровали, сеть им обучилась и она нам сама сгенерирует тесты, проверочные задания и все остальное». Но без верификации человеком это недопустимо.

Если у нас есть чёткие регламенты, инструкции, учебные материалы — мы можем сформировать базу знаний по классическому принципу экспертных систем: вручную или с помощью вспомогательных инструментов, но под контролем специалиста. Эта база будет верифицирована на полноту и непротиворечивость. Затем поверх неё могут работать ассистенты — например, языковые модели, — которые будут обращаться к этой базе знаний при формировании своих заключений.

Таким образом, нейросети не нагружаются задачей создания базы с нуля на основе каких-нибудь учебных пособий. Мы даём им готовую структурированную базу знаний, формируемую и проверяемую экспертами.

Константин Александрович Петров:

Я думаю, что экспертная система она не для того, что Роман Сергеевич сказал. Потому что миллионы правил из учебников, формализовать в живую...

Павел Романович Варшавский:

Почему? Никаких проблем нет. Если миллион таких правил есть.

Сергей Викторович Вишняков:

Смотрите, давайте представим задачу, которую описал Роман Сергеевич — генерация новой топологии микросхемы. Там есть техническое задание...

Роман Сергеевич Куликов:

Да, как правило, техническое задание на такую микросхему описано на верхнем уровне: определены размеры, энергопотребление, техпроцесс, функциональность.

Сергей Викторович Вишняков:

Хорошо, далее вступают в силу различные правила: дорожки не должны пересекаться, между ними должны соблюдаться минимальные расстояния и т.п. Более того, у нас, наверное, есть модель, есть комплексы для моделирования...

Константин Александрович Петров:

Да, конечно. Есть формализованные правила для этих комплексов моделирования...

Роман Сергеевич Куликов:

Там, как сейчас на практике, они такие сложные, что посчитать их не получается. Вот, какой-нибудь каскад усилителя можно попробовать посчитать, там зная параметры. А цифровую схему с ее параметрами утечек и задержек посчитать не представляется возможным, потому что там очень много узлов. Ее набирают в виде модели, дальше физически моделируют задержки с учетом длин дорожек, емкостей всяких взаимных. Получают подходящий или неподходящий результат.

Сергей Викторович Вишняков:

То есть, модель в этом ключевом моменте, она не полная, она упрощенная, она не учитывает многих физических процессов. Есть правила. То есть, модель, ее можно включить в контур генерации любого решения просто как основной инструмент проверки. Работает ли функционально вот эта штука или не работает. То есть, если она на модели упрощенной не работает, значит надо проводить оптимизацию, находить какие-то другие решения, менять конфигурацию.

Роман Сергеевич Куликов:

Человек обычно методом тыка, умноженным на его опыт и интуицию, начинает двигать куски топологии, местами менять. Это поближе, это подальше...

Константин Александрович Петров:

Это правильно.

Сергей Викторович Вишняков:

Есть правила, которые как раз формулирует эксперт. Не ставьте полосочки близко, потому что при оптимизации может получаться, что в простой модели у нас все уплотнится резко. Какие-то проводники будут совсем близко между собой. Человек это понимает и в нормативной документации, в регламентах, фиксирует, что будет просто пробой между соседними элементами, который является процессом нелинейным и скорее всего не моделируется сам по себе в системе. Такая физика в ней скорее всего просто не заложена. Это правило. Не сдвигать, не делать...

Константин Александрович Петров:

Все намного хуже, коллеги, эти правила проверяются уже потом, когда сгенерирована схема, причем сгенерирована она как раз таки системой проектирования.

Нейросеть можно «сунуть» только в том смысле, как предсказатель. То есть тысячу раз мы даем результаты тысячи как бы синтезов этой топологии. Да, и тысяча первую нам может предсказать. То есть как раз сейчас эта работа ведется по наработке таких баз данных по предсказанию. Что мы сразу можем быстро предсказать и поменять. Ну я в прошлый раз про это рассказывал. Но вот зашить туда все учебники и документацию, это слишком, это невозможно. Это дорого даже для LLM, тем более для экспертной системы, где все правила должны быть прописаны в явном виде. Хотя, казалось бы, есть...

Роман Сергеевич Куликов:

Это оценка, она интуитивная или есть количественные расчеты?

Константин Александрович Петров:

Это интуитивно. Вот как расставить блоки – этому, к сожалению, в Институте не научишь...

Роман Сергеевич Куликов:

Оценка того, что заложить комплекс учебников в какой-то вид экспертной или нейросетевой модели, вот тезис, он кем-то количественно оценивался или это интуитивная оценка?

Константин Александрович Петров:

Насколько это сложно количественно?

Роман Сергеевич Куликов:

Да.

Константин Александрович Петров:

Нет, я думаю не оценивался, просто не этих учебников.

Роман Сергеевич Куликов:

Учебники есть.

Константин Александрович Петров:

Нет, такого учебника, который научит расставлять правильно блоки – его просто нет...

Роман Сергеевич Куликов:

Тогда мы делаем шаг назад. Если человека нельзя научить, он же это как-то умеет.

Константин Александрович Петров:

Да, он заканчивает институт, он имеет набор базовых знаний, потом он долго доучивается...

Роман Сергеевич Куликов:

Как он понимает, что у него получается? Он делает несколько попыток последовательных в САПре, переставляет эти самые блоки, пытается добиться нужного сочетания параметров.

Константин Александрович Петров:

Он переставляет блоки, нажимает кнопку, ждет сколько-то часов и смотрит параметры, получившиеся.

Роман Сергеевич Куликов:

Он видит параметры, там 15 параметров. Смотрит сколько получилось, допустим 14 получилось, 1 не получился. Они решают с командиром насколько критично. Если критично – он переделывает. Если не критично, то командир с заказчиком согласует. Но эту же итерацию, мы говорим, что ее можно запускать автоматически.

Константин Александрович Петров:

Да, итеративно.

Можно, это как раз и делается. Можно сюда «посадить» оптиум какой-нибудь для оценки. Мы так и делаем.

Роман Сергеевич Куликов:

Можно да, но сам учебник, который описывает, что если ты делаешь классную микросхему, то тебе нужно реализовать вот, что такое канал, что такое перемножитель, этому же можно научить. Это же в учебнике формализовать можно.

Константин Александрович Петров:

Вы говорите об архитектуре системы, а здесь речь идет о конкретном блоке. Человек, который делает эту топологию, синтезирует ее, ему вообще не важно, что там внутри. Перемножитель тензорный какой-то или просто контроллер PCI-Express, вообще. Это совершенно друг с друга не связанные вещи.

Роман Сергеевич Куликов:

Думаю, что, если я правильно услышал, был тезис, что заложить в какую-то либо формализацию, либо в процесс обучения нейросети ту базу знаний, которая содержится в учебной литературе, это слишком долго?

Константин Александрович Петров:

Нет, это не слишком долго, это просто бессмысленно, потому что те знания, которые сейчас используют разработчик на месте – их в учебниках нет.

Роман Сергеевич Куликов:

А откуда эти знания берутся?

Константин Александрович Петров:

Эти знания берутся из результатов моделирования, результатов предыдущих моделирований на конкретном техпроцессе.

Роман Сергеевич Куликов:

Разработчик может описывать эти знания, они описываемые или это интуиция?

Константин Александрович Петров:

Это интуиция.

Роман Сергеевич Куликов:

Интуиция, в ней есть смысловая часть? Нейросеть может ее имитировать или нет?

Константин Александрович Петров:

Вот прямо сейчас, в эту минуту проводятся эти исследования по всему миру и в стране тоже, как раз таки и у нас, мы как раз таки к этому движемся, что вот эту интуицию постараться заменить. Потому что выработать набор правил как раз таки для экспертных таблиц не получается. Слишком сложно, слишком много. И это как раз те задачи, где используется искусственный интеллект, который надо сделать.

Павел Романович Варшавский:

Если строить правильно цепочку, то, о чем вы говорите, то это давно уже известно. Вот Александр Павлович добавит, нейросеть она что-то генерирует на основе сырых неподготовленных данных, даже может так самообучиться, даже начнет нам выдавать правильное решение. По сути, в дальнейшем нам надо оттуда вынуть экспертные знания в виде вот этой цепочки правил. Если нам это удастся, это получается обобщение вот этих знаний и получение конкретного закона или правила.

Роман Сергеевич Куликов:

Можно ли, путем обработки нейросети большого количества эмпирических данных, получить из них новые правила?

Павел Романович Варшавский:

Скорее получается так. Мы видим, что нейросеть какую-то нелинейную совершенно сложную зависимость обнаружила. Обнаружила в каком плане? Она что-то нашла. Понимаете, то, что не очевидно и не вынимается легко даже экспертам.

Роман Сергеевич Куликов:

Но мы не знаем, что из этого является причиной, а что следствием.

Павел Романович Варшавский:

Почему самое сложное для нейросетей это правильно разметить, подобрать репрезентативную выборку для обучения? Потому что, если вы это уже знаете, чему нейросеть должна обучиться, то это классическое обучение с учителем. Вот условно, у нас есть база правил из 12 или 1012 штук, и мы хотим нейросеть им обучить, такое тоже делается, то да, мы это четко знаем. Вот это она на вход получает и должна научиться с определенной точностью, погрешностью. Но есть масса задач, с которыми вот сейчас мы связываемся. Есть какие-нибудь оперативные базы данных, где все транзакции записаны в техническом объекте или в другом, но обычный человек просто не полезет разбираться. Это вот большие данные. И вот нейросети для чего? Нейросети «полезли» туда, куда человек, эксперт, оператор никогда бы. Даже если «полезет», скорее всего ему будет сложно найти там какую-то закономерность. Нейросети иногда нащупывают эти закономерности неочевидные.

Роман Сергеевич Куликов:

Не закономерности, а корреляции.

Павел Романович Варшавский:

Закономерности, корреляции, не важно. Мы видим, что нейросеть, получая на вход что-то, что мы не знаем, что должно быть, на выходе выдает правильный прогноз. Потом мы видим по факту, то есть мы это только вот эмпирически посмотрели. Посмотрите, в одном случае она нам правильный прогноз дала, ну просто по факту мы это видим, во втором, в третьем. Значит нейроструктура чему-то там научилась. Что это? Как объяснить вот то, о чем мы говорим? Объяснительный компонент. Это головная боль.

Но если мы это сможем сделать в процессе... Представьте себе, что, если специалистов предметников вы берете, например, вот кто занимается болезнями глаз, там Александра Павловича такой пример тоже есть, врачи они понимают, как должно быть. И видят по тем же графикам ЭРГ (электроретинограмма), если дойдем сегодня. То есть, по сути, если есть специалисты, если вообще бывает ситуация, когда вот нейрофизиологи, они эту нейроструктуру строят не по причине, а вот случайно давайте набросаем какую-то. А они ее строят по аналогии с биологическими нейроструктурами, то выясняется, что там есть совершенно очевидные зависимости и даже нейроструктура есть, в которых все-таки по ходу, так сказать, движения импульса можно определить, что она действует правильно.

Другой вопрос, нас-то больше интересует, когда нейросеть идет нам в помощь, что она берется за какую-то задачу, которая для нас не очевидна. Можем ли мы потом этот опыт превратить в правило или как-то обобщить, найти закономерность? Тут вопрос только в том, что если мы будем часто замечать, что вот на такие входные данные мы получаем такие выходные, то может быть и мы эту корреляцию увидим.

Роман Сергеевич Куликов:

По выявленным корреляциям удалось выявить ранее неизвестные правила?

Александр Павлович Еремеев:

Коллеги, вот можно к этому вопросу добавить? Все-таки надо исходить из задачи. Обычно методы искусственного интеллекта подключаются, когда задача «сваливается» в неопределенность. Если задача достаточно определена, работает классический метод оптимизации, то есть математическое программирование, линейное, нелинейное, дискретное. Если есть неопределенность, подключаем метод искусственного интеллекта. Опять, надо исходить из того, под какую задачу, какой метод. Вот если задача, скажем, связана с поисковым проектированием, то есть мы хотим построить новый объект. Ясно, что у нас есть уже определенные наработки в этом плане, ну или то, что говорят база прецедента. Когда методы искусственного интеллекта могут помочь выбрать наиболее подходящие методы, то есть сократить, таким образом, работу проектировщика, доставив наиболее подходящие модели, которые можно использовать в дальнейшем. Здесь не обязательно нейронные сети. Здесь во многом работают методы, базирующиеся на анализе прецедентов, аналогии,

о которых Павел Романович должен был говорить, пока не отказал. Нейронные сети. Где нужны нейронные сети? Скажем, задача распознавания, задача диагностики. Когда сеть анализирует большие наборы данных, позволяя сформулировать гипотезу, по чему создаются эти данные. Например, пульс-контроль или диагноз соответствующий, как мы с медиками наработали, это диагностика сложных патологий зрения, когда высока степень неопределенности. Патологии пересекаются. Но для этого нужно обеспечить:

- а) корректную выборку, причем большой выборку,
- б) корректную, хорошую, тестируемую выборку, на который тестируется задача.

И все равно, ну это не критически опасно, сеть может впасть в режим неустойчивости, так сказать, то есть выдать неверный ответ. Но здесь нейронная сеть, конечно, нужна, ну скажем, тот же GPT-4, там она подключена к обучающим выборкам, там, чуть ли не терабайтные объемы, так далее. Но китайцы пошли немножко хитрее – они, как и человек, сначала определяют, к какому вопросу относятся задачи, а потом подключают ту сеть, ну или, как сейчас модно говорить, агенты, которые заточены под эту подзадачу.

Да, нейронная сеть обучается, она проверяется на тестируемой выборке, после этого она работает или как помощник в диагностике, причем выдать может несколько диагнозов с разной степенью уверенности, может работать фейс-контроль и так далее. Вот здесь нейронная сеть очень будет полезна, но опять минус, нет объяснительной компоненты. В тех задачах, когда мы можем накопить уже что-то, что нам поможет решить, те же экспертные знания, те же прецеденты, которые были на объектах и так далее, вполне можно использовать другие методы, более быстрые и имеющие объяснительные компоненты. Вот это такое замечание, что не надо везде говорить, что вот бум, нейронные сети, если на нейронной сети...

Я бы на этом AI Jounpue и задавал вопросы, когда ребята говорят, у нас нейронная сеть обучена на двухсот примерах, ну это, конечно, может быть такая задача, но это несерьезно, конечно, обычно. Нейронным сетям нужны как минимум тысячи, десятки тысяч примеров, если это действительно серьезная система, скажем тот же фейс-контроль или, скажем, беспилотники и так далее. То есть максимально система должна обучиться с тем, чтобы она была надежна, скажем, хотя бы 0.9, 0.8 и так далее.

Павел Романович Варшавский:

Коллеги, я быстро. Роман Сергеевич, вопрос был: есть ли примеры? (Куликов: По выявленным корреляциям удалось выявить ранее неизвестные правила?). Есть, берем нейронные сети, которые решают задачу кластеризации, ну, маркетинговые исследования, всех, так сказать, кто посещает у нас там гипермаркет, собрали по ним данные. Не анализирует человек, не делает никаких разметок, все на вход в нейроструктуру. Да, она кластеризирует, выдает нам кластеры, а дальше мы смотрим, опа, точно, смотри, есть те, кто приходит, покупает буханку хлеба и коньяк. Ну, смотрим, что за портрет этого покупателя, или кто-то приходит, ну, то есть, я имею в виду, что интерпретацию дает человек, но, с другой стороны, мы можем на основе того, что сделает нейроструктура, увидеть, что действительно вот эти, как бы, шаблоны

покупателя, они присутствуют, и их можно делить, и видим, кого больше, кого меньше. То есть, в принципе, интерпретируем мы, но интерпретируем мы не нейросетью, а все равно экспертам, кто этим занимается, ну, вот, есть секреты.

Олег Вадимович Глухов:

Получается, задача снижения размерности входных данных, по сути, даже если этот пример представить.

Константин Александрович Петров:

Что нам мешает довести до конца и сказать, нейросеть, заодно предложи нам, как улучшить покупательский опыт, значит, ей надо ответить, что больше хлеба - больше коньяка.

Павел Романович Варшавский:

Нет, ну, я еще раз говорю, вопрос в другом, получили мы эти кластеры. Например, нам для извлечения прибыли важнее вот этот небольшой кластер, да, вот, людей, которые приходят и тратят огромные суммы на небольшие покупки.

Роман Сергеевич Куликов:

Тоже, если, например, покупателя посмотреть по регионам, то мы увидим особые портреты, например, в южных регионах, и мы не знаем, что это значит.

Павел Романович Варшавский:

Я имею в виду, что нейросеть может нащупать вот эти корреляции, проинтерпретировать мы можем их, а уж как дальше использовать эту информацию, то есть мы ее можем, там, зашивать в систему и дальше использовать или давать следующим каскадам.

Роман Сергеевич Куликов:

Получается есть прецеденты, когда нейросеть используется для поиска каких-то корреляций, но из корреляции в закономерности переводит человек. Нейросеть находит корреляцию, человек может перевести язык, если получится, если удастся.

Олег Глухов:

Тут две задачи интерпретации получается: интерпретация работы внутри нейросети и интерпретация результатов. Коллеги, у меня вопрос. Мы сейчас начали в рамках рабочей группы при нашем институте значит работу по созданию ИИ ассистента, который должен помочь программисту на ПЛИС в решении его задач. И вот вопрос какой? С точки зрения методологии или методики работы, когда мы берем, например, его рабочий процесс выстраиваем в виде дорожной карты и декомпозируем, соответственно пытаемся не всю задачу моментом за счет нейросети решить, а вот именно разбивая их на последовательные блоки это является как раз таки внедрение вот этих экспертных правил?

Павел Романович Варшавский:

Когда сеть начинает хорошо решать задачу, она же видимо находит это правило, эту корреляцию, эту зависимость и тогда она начинает его хорошо решать. По поводу генерации, давайте честно скажем, вот Сергей Викторович не даст соврать, у нас сейчас автогенерации программного кода для простеньких задач происходит сетями. Нас никто не спрашивал, и мы эту задачу себе не ставили, по сути, но сегодня вы посмотрите студенты они же активно используют. Я думаю, что вот вопрос, который звучал - и это возможно. Вопрос, что туда надо будет компонентно дать...

Олег Глухов:

Вопрос немножко другой: именно разбиение на последовательные блоки решение, это является внедрением в данном случае некоего экспертного правила? Потому что, чтобы это разбить, нужно понимать специфику задачи.

Павел Романович Варшавский:

На самом деле вот мы пока медленно движемся по кейсам, но суть в чем, ведь подходы какие есть? Если есть эксперт, который готов сразу сформулировать правила или есть в учебнике: написан там закон Ома, да, уже, то нам не надо придумывать его и ждать, когда нейросеть сама до него... Ее озарит, что вот эта зависимость присутствует. Мы можем это сразу ей по сути или в виде правила показать «вот смотри, все делай там, генерируй гипотезу, но вот это должно выполняться, да, там набор каких-то там сотни две, сотни правил». Если они у вас уже есть, зачем ее, скажем так... Ведь сеть учится на примерах. Если у вас примеры говорят о том, что это правило работает - без вопросов.

Олег Глухов:

Хорошо, давайте еще более общий вопрос задам. Если мы в целом в работу вот этого алгоритма внедряем какую-то априорную информацию - по сути это уже появление неких экспертных правил?

Павел Романович Варшавский:

Давайте я немножко сдвинусь, мы примерно об этом и говорим, Александр Павлович меня уже сегодня тоже пенал, что я еще об этом не сказал ни слова. Смотрите, если говорить про прецеденты и аналогии — это, по сути, вот некоторые переходы, о которым мы сейчас говорим. Под прецедентом можем понимать много чего, то есть это могут быть и просто там скажем некоторые снимки параметров какого-нибудь объекта, может быть в параметрической форме, может быть это какая-то структура в виде сети, фрейма там, антологии, как угодно мы это можем назвать. Основное что? Речь идет о чем? Когда это востребовано?

Когда есть набор правил и все известно по проблемной области — это здорово, вам не надо ничего придумывать, вот то, о чем сейчас говорим: зачем изобретать примеры? Предположим, проблемная область мало изучена, нет пока правил, не сформулированы. Все пандемию помнят, да, там мы же сидели ждали пока появится

штамп, там будет понятно, как лечить, от чего или мы накопим опыт относительно того, как правильно лечить там и так далее. То есть был вот период ожидания, когда у нас есть только какие-то частные отдельные примеры. Вот то что называется прецедентом, да, или аналогом. То есть, по сути, у нас есть всего один пример. Но с индуктивным подходом все более или менее знакомы, когда мы действительно, вот, от этих частных пытаемся перейти к общему заключению. Аналогии и прецеденты – попытка перейти к заключению на основе всего одного единственного примера, случай какой-то, инцидент исключительный, пока вот мы его в тираж не запустили. Да, вот он один такой случился, что делать здесь? Нужно поискать какое-то сходство аналога с чем-то, что имело место быть, то что о чем идет речь с прецедентом. Понимание для нейроструктур – это может быть обучающий пример. То есть, вот, у нас где-то на какой-то станции аварийная ситуация возникла из-за этого. Теперь, конечно, логично этот пример инцидент по остальным объектам подобного типа, да, тоже его иметь в виду, что такое может случиться. Как сегодня там про Фукусиму и все остальные можно говорить. По сути, вот у нас была работа, которая была связана с чем? Коллеги разрабатывают специальную систему «скады и мониторинг процессов» на атомном энергоблоке. У них там есть все регламентированные правила, включая штатный, нештатный, даже высадка инопланетян. Все написано, что надо делать на атомной станции. Мои попытки объяснить, что можно использовать подобный инструмент для более гибкой работы долгое время не воспринимался коллегами, потому что они говорили...

Роман Сергеевич Куликов:

На самом деле в человеческой практике это используется, после каждого инцидента в отрасли идет разбор полетов...

Павел Романович Варшавский:

Это когда инцидент уже произошел. Вопрос: а можно ли не доводить до инцидента и поймать ситуацию раньше? Вот сошлись мы с НИИКА в то время на том, что да, у вас есть регламент, есть своя база знаний, вот все экспертами – никаких нейросетей. Они «гоняют» эту базу знаний в реальном времени. Если много параметров разбиты на подсистемы, и по сути, то о чем мы как раз и договорились, и они поняли, что это может давать им, как интерес - внедрение некоторого компонента. Это вот как раз такую методику мы и сделали. Вот часть, структурированная в то время, но сегодня у нас любят это в виде антологии называть, это семантическая сеть, описывающая ситуации на системе охлаждения зоны реактора, то есть одной из подсистем. Где параметрические позиции есть, есть описание ситуации, то есть что-то по параметрам и есть некие рекомендации. Вот эта ситуация она скажем так менее изучены, то есть, где мы наблюдаемую текущую ситуацию, когда там задвижки всевозможные. Там да показатели, описано да, и есть более полноописанные ситуации, там, например, третья, да, там вторая. Где мы можем уже видеть даже и рекомендации, которые даются оператору в случае, если наблюдается соответствующая ситуация. Но здесь вот просто действительно параметры структурированы, тут получается в виде даже, пока

вот, такой не столько сети, сколько дерева, но тоже и сетевая структура здесь вполне может быть.

В основном зачем это сделано? Для того, чтобы сравнивая сходство ситуации, выявлять некое подобие и на основе этого подобия предлагать возможное решение. Вот как раз для тех единичных примеров. Никогда у нас правил кем-то записано, а когда мы просто видим сходство, что вот одна ситуация очень похожа на другую и возможно придется принять то или иное решение. Соответственно на объекте, что мы делали по подсистемам атомного энергоблока, по сути, была дополнена база прецедентов, порой их как инциденты воспринимают. То есть это была задачка, когда эксперты какие-то штатные примеры давали, их не так много кстати. Вот здесь самое интересное, что не нужно раздувать эту базу прецедентов накопленного опыта. То есть это не стоит задачи весь регламент превратить в прецедент. Такое можно сделать, но тоже пришли к выводу, что это не так важно, важнее дать прецеденты, которые существенны. Пример, аварийные ситуации, которые вызывают, инциденты, которые по отрасли рассылаются. То, что нужно подсветить оператору. Ну, вот на компенсаторе объема мы поэкспериментировали. В каком плане? У коллег с НИИКА была соответственно система, которая по регламенту работает, то есть система мониторинга и в оперативной базе всю информацию дает. И мы поняли, что на эффективное применение возможно в параллель, то есть система работает как? Есть их решатель классический, который крутит их базу, то есть к ним мы не предстаем. Соответственно есть вот эта система, которая через прецеденты успевает в определенный момент, там, не каждые, например, 3 секунды, через какой-то регламент пытается определить текущие ситуации: напоминает что-то из прецедентов, и если напоминает, то с определенной степенью сходства дает эту информацию дополнительно оператору. Зачем это нужно? Но, по сути, действительно человек способен не ждать, когда у него уставки сработают. Да, вот тут пример создания этих прецедентов, а вот как раз работает их система, до которой по оперативной базе, как бы сейчас мы говорили цифровой двойник, в то время еще это не так называлось, то есть диагнозы/прогнозы выдает система своя у них, которая всю базу норматива, скажем так, имеет. И, соответственно, вот здесь еще появляются рекомендации в виде вот таких соответствующих ситуаций.

Где мы что видим? Было видно несколько моментов. Во-первых, как на объекте атомы нагреваются. Подходим к тому, что надо сейчас будет задвижку закрыть, чтобы у нас температура снизилась. У них стоят куча уставок. Причем не 2, не 3, а много. Вот он приближается, приближается, потом дальше срабатывает и сбрасывает. Но прецедент при этом в какой-то момент пишет, что вы близки, там вот уже сейчас придется что-то делать. Да, то есть основная идея какая? Оператор опытный, конечно, ему может быть такая подсказка не очень нужна. Хотя, с другой стороны, если он видит какой-то прецедент или инцидент, с которым степень сходства становится... Фактически что происходит? Мы договорились о чем? Вы понимаете, что у вас оператор за большее количество времени уже увидит информацию о том, что он приближается к какой-то нештатной ситуации. И он не будет себе читать, когда все замигает и заморгает. Он будет только кнопку нажимать. А то он может, если там сидит, скучает и вдруг увидел,

что у него что-то близко к какому-то инциденту, он может начать какие-то предпринимать действия. Если он опытный оператор, он может обратить внимание на какую-нибудь еще схему под системой. Если это стандартно, то дальше будет спать... Вот здесь, по сути, и вот основное, что мы делали. У меня и ребята в магистратуре, операторы тоже крутили такие вещи. Прецеденты. Их немного. Они должны быть эксклюзивные для того, чтобы на них реагировать. И самое важное, что были попытки как раз нейросети. Вот есть у тебя база прецедентов. И посмотри, как работают нейросети и прецеденты. С прецедентами в чем плюс? Прецедент это то, что мы уже знаем. То есть это, по сути, пример, для которого все понятно. И мы нашли на него похожие. Слушайте, вот это похоже. И у нас легко. Сразу идет обосновательный компонент. А с точки зрения нейросети, представьте себе, что если нейросеть учится на этих шаблонных примерах и тоже их соответственно хорошо знает, то значит и нейросеть будет выдавать тот же самый, с той же, скажем так, точностью вам результат. То есть вполне может быть именно как очень репрезентативная, скажем так, хорошо размеченная база этих примеров для обучения нейроструктур. То есть вот такой вариант тоже взаимодействия возможен.

Сергей Викторов Вишняков:

Вот как раз же и вопрос, прецеденты-то они могут быть редкими. То есть для базы, которая используется для обучения нейросети на самом деле, либо мы их будем искусственно отчеркивать.

Павел Романович Варшавский:

Коллеги, тут и получается придумать тоже прецедент. Это не панацея. Если у вас пример какой-то появляется, новый, вот то, что мы говорили, его всем передают и дальше мы уже можем учитывать, что такое может иметь место быть. Соответственно понятно, что если просто использовать практику прецедентов и накапливать опыт больше и больше, то мы просто будем накапливать много однотипных ситуаций, которые будут только подтверждать какой-то факт. На самом деле с прецедентами проще от прецедента перейти к обобщению и получить базу правил. Это вообще легко. Вот те же уставки и так далее. То есть если мы видим, что у нас там сотни прецедентов, это вот классические дальше индуктивные схемы, можно перейти к обобщению и получить, что вот у нас, например тысяча прецедентов об одном и том же. Нужен реально один, который как раз эту зависимость и показывает. Из него мы можем перейти к правилам. Здесь появляется обосновательный компонент. Это уже элементы, которые бывают полезны. Как это внедрять в учебные процессы? Безусловно аналогии и прецеденты — это тоже, как еще и Минский говорил, это тот инструмент, который тоже в выдвижении гипотез мы опять вот туда же попадаем. То есть на основе одного примера мы там пытаемся сделать предположение, что это похоже на другой пример. И давайте попробуем лечить пациента также. Но это не значит, что достоверность. Примерно так одного вылечили, а у другого летальный исход. К сожалению, когда идет индуктивная схема на основе одного примера, риски

тоже и степень уверенности примерно такая. Но что еще вам показать, чтобы вас не замучить совсем?

Константин Александрович Петров:

Скажите, пожалуйста, все-таки, кто определяет похожесть или непохожесть инцидента? Кто понимает, что вот это инцидент?

Павел Романович Варшавский:

Основная идея алгоритмов, которые строятся по прецедентам, это вопрос, какую метрику мы вводим для определения сходства. Если это параметрическая ситуация, как вот здесь в примере, который мы рассматриваем, где не нужно структуру смотреть, хотя здесь и структурированные алгоритмы тоже есть. Основная идея простая. Мы берем, например задачу простейшую параметрическую смотрим. У нас там, например сотни параметров. Если у нас в этой сотне параметров мы метрически определили важность параметров, например там давление, температура играет большую роль, определили набор параметров. Дальше мы используем классические метрики, хотите Евклидову, хотите другие метрики. Для качественных оценок можно применять при разработке системы.

Константин Александрович Петров:

Кто эти «мы»?

Павел Романович Варшавский:

При разработке системы мы с вами в алгоритм определения степени сходства закладываем определенные метрики.

Константин Александрович Петров:

Заранее? (отвечают: да) Это определяет границы, значит вот это такие прецеденты.

Павел Романович Варшавский:

Нет, определяем метрику. Как мы будем считать показатель сходства?

Сергей Викторович Вишняков:

Разность температуры квадратных параметров...

Роман Сергеевич Куликов:

Если говорить про топологию, это некоторая безразмерная метрика разницы между полученными характеристиками и требуемыми.

Константин Александрович Петров:

Но при этом характеристик может быть сотни. (отвечают: вполне) Хорошо, мы их с разными коэффициентами используем.

Павел Романович Варшавский:

Нет, эта задача не сложна с точки зрения вычислительной части, она понятна в параметрах.

Константин Александрович Петров:

Вопрос ответственности, кто вот это вот решил? (отвечают: генеральный конструктор решил)

Павел Романович Варшавский:

Ответственность на программы, места, размеры, диапазоны, плюс ко всему и, наверное, в вопросе, когда определяется сходство, задается пороговое значение. В частности, на АЭС мы вводили порог. Какой смысл выдавать прецеденты, у которых степень сходства ниже 50 процентов или даже 70–80. То есть мы можем говорить там, где мы это видим. С прецедентом они у нас сегодня могут быть разные. Это может быть рисунок, график. Они похожи или нет. Метрики могут быть различные. Вы правильно говорите, кто будет определить? Конечно, при постановке мы должны в предмет области определиться. Плюс ко всему строили варианты, когда оператору выдается решение по нескольким метрикам. То есть вы видите в одной метрике вот такие прецеденты, в этой другие, третьи. И дальше только эксперт по предметной области может сказать, что это в большей степени правильно или нет. По итогу то же самое, как с интерпретацией результатов от нейросетей.

Сергей Викторович Вишняков:

Еще можно же просто даже тесты провести. Есть какие-то известные решения. Вы можете посмотреть, какие метрики дали более высокие проценты.

Павел Романович Варшавский:

Самое главное, чтобы снять вопрос, надо не забывать, что прецеденты и аналоги — это инструменты выдвижения гипотез. И по сути метод правдоподобный. То есть это не попытка вам как в дедуктивной схеме привести, что 100% так. Это решение, которое он предлагает. Вот это не то, что должен делать оператор. Это лишь предупреждение, рекомендация. Достаточно, что при 80 степенях сходства ему надо бежать и проверять, что уставки так. И надо что-то менять. Это как раз дополнительная информация, не относящаяся к достоверному заключению. Мы не прогноз и диагноз ставим на основе этого. А скорее это инструмент, еще раз, выдвижения некоторой гипотезы, которая может быть полезна для недоведения...

Роман Сергеевич Куликов:

Ситуация на 80% напоминает Чернобыльскую (это была шутка).

Павел Романович Варшавский:

Сразу уже передумает пить. На самом деле, я могу сказать еще один момент. Мы до этого не дошли по внедрению. Часто действительно в случае каких-то уже аварийных ситуаций. Разбор полетов. И вот здесь мы можем сами генерировать возможные

причины. А с другой стороны, мы можем смотреть. Есть количественная оценка, что у вас, ребята, вы можете искать, но у вас очень похоже вот на это и на это.

Константин Александрович Петров:

Это лишний повод. У вас похоже на это на 70, а на это на 80. И оператор должен был понять, что 80 — это больше, чем 70. Следить за еще одним ноутбуком, помимо лампочки.

Павел Романович Варшавский:

Я могу сказать, что бывают ситуации, когда 100% сходство. То есть полностью идентичная ситуация. Но почему нет? Если у вас все параметры сложились один в один. То есть у вас будет также прецедент со стопроцентным совпадением.

Александр Павлович Еремеев:

Про нейронную сеть. Пример с медиками. Нейронную сеть обучали распознавать диагнозы. Диагнозы там довольно запутаны. Потому что накладывается... нарушение сетчатки, диабет и так далее. Вот нейронная сеть была обучена. Ну, использовали глубокую нейронную сеть, каскадную. После этого, вот этот вариант мы предлагали медикам. И медики фиксировали, да, действительно правильная сеть. И подтверждали, что по этому набору данных, действительно, диабетическая ретинопатия. Но есть в какой-то степени неуверенность. Или неправильно. Вот здесь, практически, нейронная сеть была очень полезна. Но, опять-таки, полезна, когда достаточно серьезные выборки. Причем, опять-таки, нейронная сеть. Вот здесь, да, если система не могла распознать или распознавала неправильно, то эксперт подключался, скажем, какой дополнительный параметр надо ввести или какой вес повысить, какой понизить. Вот здесь пример использования нейронной сети. Правда для того, чтобы подготовить корректно обучающую выборку, пришлось предварительно использовать, так называемый, препроцессинг. То есть, отбрасывать второстепенную информацию, чтобы не «захламлять», как говорится, нейронную сеть, и оставлять наиболее существенную. То есть, там подключали мы и вейвлет-преобразования, так сказать, и сжатие информации и так далее. Вот в этом плане нейронная сеть была очень полезна. Но, опять-таки, все зависит от обучающей выборки, так сказать. Насколько она надежна, насколько она корректна. И если, действительно, накоплена хорошая база прецедентов, то эти прецеденты можно делать как обучающий выбор. То есть, систему можно сгенерировать гипотезу достаточно правдоподобную. Если этой обучающей выборки не будет, или, так сказать, эксклюзивный опыт, то здесь, наверное, нулевой или даже отрицательный результат будет.

Олег Вадимович Глухов:

Понятно. Спасибо, Александр Павлович. Павел Романович, я предлагаю немножко нам ускориться. Нам презентацию-то разрешите всем участникам предоставить?

Роман Сергеевич Куликов:

Публиковать ее можно на разделе портала?

Павел Романович Варшавский:

Почему бы и нет. Коллеги, ну вот Александр Павлович уже к медицинской тематике перешел. Действительно, с этой тематикой и с Институтом Гельмгольца мы довольно плотно работали. Причем по нескольким направлениям. И с методами ИИ работали. И вот кафедра УИТ: у нас как раз там в свое время Александр Ильич Колосов занимался. Они там с дифференциальными уравнениями решали в приближениях с этой части. Ну и вот, Александр Павлович и его группа работала как раз с нейросетевыми методами использования когнитивных образов. Ну вот то, что там по глазам показывали. Вот здесь, что можно отметить, что здесь действительно нейроструктуры, они реальные. Биологические. То есть в глазу там есть реально нейроструктуры. Об этом речь. Когда мы смотрим на модели реальную и ее моделируем, то конечно здесь природа за многое, очень часто эффективное решение там и получается. Действительно здесь основная беда была в чем? Медики сейчас они работают как у них там американские или любые другие приборы, которые на вспышки глаз реагируют. Они снимают показания и тебе распечаточку. И там уже врач специалист смотрит и постановку диагноза делает. Наши там попытки поработать, там отправлять студентов даже, чтобы у них выборка появлялась. По набору посмотреть что и как. С 15 человек три здорового глаза попадают иногда в лучшем случае. А то и не одного. Вот. Основная беда у них — вот, например, здесь прецеденты или примеры обучающие. Просто сложно найти здоровый глаз, чтобы вот реально его тестить и по нему получить какую-то исходную информацию. Соответственно вот здесь, в этом контексте на самом деле задачу, которую решали в группе у Александра Павловича, то есть, по сути, что делали? На вход поступала вот такая вот картинка с этими альфа- и бета-волнами. Соответственно распознавали, не строили эту модель приближения, как коллеги с УИТ. А здесь действительно просто график. Рисунок. Изображение нейросеть получала разные волны. Обучалась и знала, что вот эта волна — это некая схематизация звучания. По сути, было как? Убирались шумы, убирались те вот волны, которые не играют принципиальные роли для постановки диагноза и дальше уже имея шаблонны нейросеть на этих примерах обучалась, дальше уже может...

Роман Сергеевич Куликов:

Изначально, форма вот этой кривой, ее по опыту нарисовали врачи или откуда она взялась?

Павел Романович Варшавский:

Нет, у них приборы... Действительно кривая у каждого будет своя, но у этих 15 не всегда здоровый глаз. Ее выдает прибор серии вот этих вспышек, когда вы смотрите глазами.

Роман Сергеевич Куликов:

Если я правильно уловил, то диагностический вот этот комплекс, получив кривую конкретного пациента, что-то в ней пытается найти, но при этом он...

Павел Романович Варшавский:

Нет, он ничего не пытается найти, он вас просто снимает показания, реакцию сетчатки на вспышки. Кратковременные и так далее.

Роман Сергеевич Куликов:

Какой-то референс?

Павел Романович Варшавский:

Нет, он вам просто выдает, просто вместо кардиограммы, он вам выдает вот этот график того, как реагирует сетчатка на вспышки.

Александр Павлович Еремеев:

Можно я тогда поясню? Мы с ней они работали. Это прибор у некоторых ... который в определенной степени это кардиограмма глаза, но только не электрические сигналы, а вспышки. И он выдает вот эту электроретинограмму, проинтерпретировать которую может только специалист физиолога. А текущий врач он ее не может проинтерпретировать, потому что это довольно сложно. Здесь альфа-волна, бета-волна наиболее информативные, ... не информативные. Наша задача была, к сожалению, прибор немецкий или японский, исходники закрыты. Одна из задач была снять эту информацию, как говорится, цифровать непрерывные сигналы в дискретную И второе, все-таки построить определенный классификатор, который направлен на электроретинограммы, может в какой-то степени точности классифицировать ситуацию и помочь при диагностике. Диагностика там довольно запутанная. Это была задача...

Роман Сергеевич Куликов:

Исходный базис для классификации откуда брался? ...

Александр Павлович Еремеев:

К сожалению, мы не могли получить выбор по нормальным глазам, чтобы у нас было нормальный образ жизни. Уже когда зрение потеряно, ну и, к сожалению, по студентам, мы две группы полностью предоставляли. Ну, там, например, несколько здоровых классов оказалось всего. Поэтому обучить сеть, чтобы она сравнивалась с нормальным зрением, это было довольно проблематично. Вот почему нейронная сеть и не давала хорошей классификации, в принципе, это было принятым. Но, опять-таки, обучающая выборка на наш взгляд была еще недостаточна, чтобы на этом можно было давать оптическую рекомендацию. Вот что я хотел добавить.

Павел Романович Варшавский:

Я попробую пояснить. Задача строилась так. Учитывая, что у коллег-медиков никаких наборов данных реально нет, есть только конкретные пациенты с вот этими снятыми показаниями. То есть у них вот эти графики, эмпирически полученные графики, причем они их не сильно систематизировали. Соответственно, что пришлось делать? Нам пришлось собирать вот эти по итогу распечатки, условно говоря. И врач-офтальмолог, он говорит, да, вот это диабетическая ретинопатия, это такая стадия, это такая, а это здоровый глаз один попался. Вот нам пришлось разметить, так сказать, выборку вот этих примеров. И дальше уже в нейросеть обучать. Обучили нейросеть, и нейросеть...

Константин Александрович Петров:

Нет, есть эксперты, разметившие, здесь как бы всё нормально.

Роман Сергеевич Куликов:

Как у вас, эта топология удачная, а эта неудачная.

Константин Александрович Петров:

Да, у нас сейчас идёт разметка.

Павел Романович Варшавский:

Это в основном как раз. Получается, что как только мы имеем наборы примеров, там, удачные, неудачные, вот с этими постановками стадий, дальше эта нейросеть, уже без врача, начали давать ей тестовые наборы, ну или новый пациент. И она осталась определённой точностью, да, как итогом, конечно, заменить врача. То есть, особенно это для ординаторов, да, то есть ещё молодого, неопытного, он не бегаёт, за халат не дёргает, врача там, это точно диабетическая ретинопатия? А у него есть подсказка в виде нейросети, мы говорим, здесь там такая-то стадия.

Константин Александрович Петров:

Вероятность?

Павел Романович Варшавский:

Да, да, да. Вы же слышали, Александра Павловича, вся проблема в чём, медики вот эту выборку, у них аппараты, которые им поставляли, да, вот это не наши, где мы можем выгружать, потому что если бы мы могли это всё оцифровывать, собирать, смотреть, как вот эти изменяются там, альфа-волна, бета-волна и так далее, то есть вот то, что и УИТ занимался, они пытались оцифровывать и расшифровывать.

Константин Александрович Петров:

То есть, у аппарата ограниченное количество использований?

Павел Романович Варшавский:

Нет, у аппарата, он выдаёт вам распечатку. Вы получаете... Сканер, предобработка. Сканер, предобработка, убираем шумы, переводим, делаем нейросеть, обучаем нейросеть, которая дальше уже, делаем наборы данных... Об этом и речь, то здесь пришлось с нуля собирать, после этого расшифровывать, и параллельно, вот то, что занимались, это и попытка всё-таки эту кривую получить не в виде картинок, а действительно, чтобы эта оцифрованная была информация полностью. Но это, если говорить дальше. А дальше действительно, то, о чём мы говорим, если говорить, как сейчас, так же снимки вот эти, скажем так, то, что сейчас идентифицируют и распознают нейросеть, не хуже врача, там, узиста, да, или там, рентгенолога, кто там оценит. Это же примерно туда же. В те же моменты, только вот за глазами. Единственное, что ещё для врачей-специалистов делали когнитивный образ в виде вот этого набора сетчатки с их, и как раз было видно, да, вот в процессе, ну, здесь я, значит, покажу концовку, то есть тоже, то есть вопрос, как всегда в предметной области люди по-разному понимают результаты. То есть, вот видите, в качестве результатов система выдавала, вот как раз, картинку вот этих слоёв сетчатки, да, при соответствующих показателях. И здесь уже...

Роман Сергеевич Куликов:
Ваша система?

Павел Романович Варшавский:

Да, да, это уже наша. То есть, вот это как раз, вот то, как мы видим вот здесь справа, нам это ни о чём не говорит, все эти, ну, да, а врачу как раз это говорит о чём-то.

Павел Юрьевич Анучин:

Павел Романович, вопрос-то, если мы в нейросеть закладываем итоги рецензирования, ну, заключения врача, да, а не то, за счёт чего врач рецензирует. Он же, наверное, какими-то правилами руководствуется, чтобы отрецензировать то или иное с ним...

Павел Романович Варшавский:

Нет, нет, он смотрит на эти волны, он смотрит на волны. (отвечают: а как?) Нет, в этом-то вся и беда...

Роман Сергеевич Куликов:

У всех специалистов сходится к тому, что в каждой предметной области есть нечто, что плохо формализуется. Ну, это никуда не едет. Но при этом специалист...

Павел Романович Варшавский:

Да, есть Иван Иванович с молотком, который знает, в какое место надо ударить.

Павел Юрьевич Анучин:

Проблема в том, что, условно, нейросеть упрётся в потолок конкретного специалиста, который для неё это размечал, а не в...

Константин Александрович Петров:

И вот хочется верить, что как раз-таки она пойдёт дальше.

Сергей Викторович Вишняков:

Первое, конечно, когда мы говорим, что плохо формализуется, мы неправы. Никто их не пытался никогда формализовать. Никому никогда не надо было. Не было компьютера. Зачем учить компьютер правилам, если нужно научить человека? Мы сейчас как раз свидетели того, как парадигма меняется. Нам компьютер надо научить.

Роман Сергеевич Куликов:

Об этом мы и говорим.

Сергей Викторович Вишняков:

Поэтому тоже учебники, они показывают вот такую картинку. Вот такая.

Павел Романович Варшавский:

Нет, это, конечно, тема для одного. Вы поймите, я эту проблему давно поднимал. Попробуйте научить писать картины гениально. Или композицию. Вопрос инсайта открыт. Как? Мы же с вами, давайте так, нейрофизиологи, кто занимается изучением мозга человека, там, как это, белых пятен, чёрных дыр масса. То есть есть процесс, который хорошо...

Роман Сергеевич Куликов:

Первый тезис, что есть принципиальное наличие областей знаний и умений, которые не подлежат формализации, допустим. Вот тезис такой. Другой тезис, Сергей Викторович сказал, что просто их не пробовали. Потому что другой был интерфейс общения, так далее, человеческий, а не компьютерный.

Сергей Викторович Вишняков:

Нет, ну что же без врача? Не проводилась же диагностика. Какой смысл формализовать, если надо врача учить?

Роман Сергеевич Куликов:

Хорошо. Вопросик такой. Ну, я предлагаю уже там на полтора, или на больше часа, начать упражняться. Есть вопрос. Выскажите, пожалуйста. Предлагаю всем, это важно для педагогики искусственного интеллекта, на мой взгляд. Вот мы сейчас по умолчанию находились в такой парадигме, когда инструменты искусственного интеллекта, там разные подходы к обучению к конкретной предметной области, в которых, так или иначе, в виде какой-то информации, знания вычленялись и по-разному использовались. Человек обучается, например, с детского садика, со школы, с

общим предметом, потом специализируется. Насколько, по-вашему, например, ну, если там сейчас позвольте такой блиц, опрос там, ну, коротко там, нужно, не нужно. Например, если мы, предположим, не имеем ограничения по времени вычислительным ресурсам, не будем пока об этом отталкиваться, то насколько улучшить ситуацию с той же интуицией или еще чем-то уже в искусственном интеллекте, в интеллектном исполнении, если, например, прежде чем погружать в инструментную область, какой-то тот или иной инструмент в искусственном интеллекте, сначала так же обучать, ну, условно, в программе школы и базовых первых, вторых, третьих курсов, прежде чем обучать специальному знанию. Это даст значение или не даст?

Константин Александрович Петров:

Конечно даст. Машина, ребенок, все, вот она и есть. Почему бы и нет? Другое дело, что какие-то знания, наконец-то же, отбросят последствием, ну и все. К тому же при неограниченных ресурсах, вот как вы ставите задачу. Ну пока, чтобы не упираться. При слабо ограниченных ресурсах хорошо, это все прекрасно. Вот, честно говоря, пока нету какого-то понимания, есть как бы человек, это как бы качественный мозг, да, как нам кажется, он может в искусство, он может в науку, он может, извините меня, в религию, но при этом количество... Да, ну, конечно, там, недостатки у машины есть, она может сломаться, конечно, там бензонасос отвалится, еще что-то, конечно, там импульсы, эпилепсия, все такое. Но он, в принципе, получается, что-то он небольшой. Потому что человек, он как бы не может держать в голове слишком много. Тем более, далеко не каждый человек, да, может держать в голове много. Там, ребенок, он может, ну, потому что он свободен от каких-то там нагрузки, он может там быстро, он тратит на обучение, да, когда он научится, там, человек там... Я там думаю одновременно о том, что какие продукты купить домой, что на работе, и вот опять-таки в тематике семинара, да. А вот насколько хорош искусственный интеллект тем, что может быть просто количественно сильно увеличен и больше сосредоточен на задаче, вот. Вот тут возможно просто для прорыва какой-то момент возможности, да. То есть там, не знаю, вот опять сосредоточиться на одной задаче, и весь мозг посвятить этому в данный конкретный момент. Плюс удержать это все в памяти, а память компьютера сильно больше, чем память человека. Это вообще, ну, замечательно. Мозг Эйнштейна был хорош тем, что он был там переразвит в определенных местах, которые позволяли просто сосредотачиваться на одной задаче, отбрасывая и забывая все остальное в моменте. Поэтому количественно искусственный интеллект может, наверное, пойти далеко вперед. И правильно его обучить, он может отбросить лишнее и взять только то, что надо. Мне так кажется

Роман Сергеевич Куликов:

Хорошо. Понятно... А Вы как считаете, имеет смысл выходить сначала за рамки предметной области, а потом только...

Павел Романович Варшавский:

Коллеги, это тоже уже известный момент, ведь базы знаний есть специализированные, предметные, а есть базы общих знаний. Это же тоже как физика всего, да, так и вопрос, мы же оперируем не только понятиями в предметной области, когда принимаем решения. Вот то, что было сказано, мы себя осознаем во времени, в пространстве, знаем много всего дополнительного, и приходя в предметную область, результаты часто получаются на меж, так сказать, моментах. Некоторые задачи, как у Перельмана, это как раз там, где надо выйти за границы предметной области, топологии было. Поэтому на самом деле это так и решается. Попытка создавать онтологии общих знаний, то есть погружать то, о чем мы говорим, тебе нужно решать там задачи в области термодинамики. Но это не значит, что не надо, чтобы в сети или где-то в системе была база общих знаний, которая, может быть, нам кажется, сейчас не нужна для принятия каких-то решений, отыскания решений в предметной области. Но в основном у людей, как раз эти базы все подключены. То есть мы порой не нужны отбросить, мы их не отбрасываем. Но если верить опять нейробиологам, то все, что мы видели в долговременной памяти, все хранится. То есть оно никуда не исчезло, оно есть. Другое дело, что у нас очень неэффективные механизмы изъятия из этой памяти, то есть если в оперативной условно говоря в кавычках мы довольно активно это используем, то все что в долговременной тут вот как разведчиков учат, да там они могут карту заполнить быстро визуально еще что-то, здесь примерно та же ситуация. То есть если говорить, наш мозг как компьютер в кавычках рассматривать, то действительно, наверное, мы это все интегрируем. Вопрос что используется в конкретных моментах там для принятия решения, для отыскания его, это вот как раз тот сложный момент. Безусловно это расширяет возможности ИИ там условно говоря, вот особенно в области педагогики. Вот сейчас уже как на рынке наплодили куча джуниоров по ИТ и говорят, а нам нужны мидлы там или синьоры, а вы тут нам джуниоры. С другой стороны, если сейчас мы выбросим этот слой джуниоров и поставим туда ИИ, он с этим скорее всего справится, но откуда будут браться мидлы и все остальные. (ответ: надо видимо учить каких-то промпт инженеров). Как только вы через что-то перепрыгиваете, это значит, что не сформируется определенный пласт вот этих знаний, либо вы их должны где-то искусственно, тогда как-то через это пропустить там в ВУЗах, колледжах или еще где-то. Если этого нет, оно вот нейроструктуру это не приобретет. Ну как мальчик Маугли, если ты не общаешься в определенном возрасте, то дальше языковые навыки и не появятся. Здесь боюсь, это уже наша специфика. Ну и второе, если про педагогику в прошлый раз нас не было, тоже вот на Пospelовских чтениях в конце прошлого года мы говорили на нашем семинаре. Проблема в том, что вот применение ИИ. Тут не вопрос надо или не надо даже, хотя кодекс применения в образовании сейчас и в Думе, и в ВУЗах обсуждается активно. Вопрос в другом, что сейчас студент пришел уже тот, который этим пользуется, и скорее мы уже должны догонять их и переходить к тому, что действительно старым методикам, да вот как обучались, сегодня просто студенты уже вот большинство из них, давайте так скажем, они, вообще говоря, уже переориентированы. Цифровое или электронное у них образование. То есть я не говорю, что их вообще нет, но их меньше. Основная масса поменялась и действительно

сегодня там с генерацией с сетями и так далее это будет. Оно никуда не деться. Есть инструмент, есть калькулятор, да вот кто там пользуется, техническая линейка или корень квадратный на листочке умеет извлекать, когда есть вот гаджеты. И сегодня скорее всего нейросети, ИИ будет давать полезный инструмент. Единственное, что вот проблема в образовании, вот мы прослеживаем, как мы давно этим занимаемся, подготовка специалистов у ИИ, надо четко разделять. Есть часть специалистов, которых надо обучить технологиями ИИ, вот то, о чем мы говорим. Помощники умеют пользоваться кнопками. Но совершенно то, что сейчас и у нас, и на Западе не обращает внимания, это на обучение специалистов именно теоретическим, фундаментальным основам ИИ. То есть те, кто могут это делать. Сейчас на хайпе идет процесс такой, что любую домработницу можно научить пользоваться ИИ. И так должно быть. И самое главное, вот идет вопрос о замене. Да, мы можем заменить, но мы же самое интересное, как говорят, пользование вот этими трансформерами и генеративным ИИ. Сегодня специалисты из областей индустриальной говорят, что это лучше не для джуниоров. Это как раз там плохо, они говорят. Они не включают мозг, не начинают ничего делать. А вот для мидл и выше, да, это полезный инструмент. Вот как объяснить студенту или тому, кто там претендует на джуниора, что это инструмент, а не то, что у них желание написать не самому программу или ее сгенерировать и потом отредактировать, сделать эффективнее. А у него просто сделал и отдал, все и забыл. Примерно это сейчас наблюдает в отрасли. И сейчас у них тоже головная боль начала по этому поводу. Они проходят здорово все эти собеседования вместе с ChatGPT, а дальше начинается проблема после испытательного срока. Почему они плачут, зачем нам столько вот этих джуниоров, которые вот только тебе из сети вываливают решения. А как их сразу перевести на мидл уровень? Ну вот пока решения они у себя не нашли. Ну а то, что касается у нас, я думаю, что мы будем из года в год все больше и больше видеть, что если сами не начнем вот то, о чем мы сейчас говорим, то скорее мы будем видеть это снизу, голосовать нам будут.

Олег Вадимович Глухов:

Это вот мы опять возвращаемся к прошлому семинару, мы поднимали вопрос методики оценки. То есть они буквально, что в образовании, что вот в отрасли эти методики, как оценивать производительность вот этого инженера или там программиста, который это использует, вот сейчас пока непонятно. То есть мы используем старые методики для того, чтобы оценить текущий результат. И вот здесь надо как раз таки делать шаги, в том числе за счет семинара это обсуждать или семинаров, выявлять проблемы и решения.

Роман Сергеевич Куликов:

Да, в том числе это возможное решение, потому что мы вроде встали на этот путь, но он такой туманный, надо осмотреться.

Олег Вадимович Глухов:

Сергей Викторович, у вас какое мнение?

Роман Сергеевич Куликов:

Да, надо ли ИИ учить физики с химией для задачи с топологией?

Сергей Викторович Вишняков:

Я не уверен, что нужно, потому что, когда мы обучаемся, у нас есть система мотивации нашей, да, мы обучаемся не ради чего-то абстрактного, да, у нас родители заставляют, или мы там хотим куда-то потом подступить, и вообще наш жизненный опыт, он, ну мягко выражаясь, он не совпадает с жизненным опытом нейросетей. Она не испытывает таких удобств, неудобств, эмоциональных каких-то поощрений или наказаний, ну то есть у нее мотивация-то другая. Мы ее заставляем изучить формально, пройти учебники, там пройти, пройти что-то. Вопрос заключается в том, она-то воспринимает это, что ну хорошо сформировались какие-то связи, какие-то веса настроились. Для нас это немножко другой опыт, на мой взгляд. Для нее это будет другой опыт. И я не очень понимаю, зачем этот опыт ей нейросетей нужен. В конце концов, хотим решить какую-то задачу, мы ей говорим, вот такая картинка соответствует ретинограмме, она говорит, ну вот, а вот на уроке химии в школе мне говорили, что так вообще, не надо делать. То есть, не совсем понятно, мне кажется, это не целесообразно, по крайней мере, на текущем опыте. Вообще, в принципе, непонятно, зачем повторять нейросети, опыт человека, когда мышление разное, цель и задача разные, мотивация, мягко говоря, есть разная.

Роман Сергеевич Куликов:

А у вас есть мнение?

Артем Анатольевич Птицин:

Да, мнение есть, но тут все зависит от того, под какие задачи делается ее стенд. То есть, допустим, для похожей задачи по выявлению каких-то аномалий в глазах, по конкретной кардиограмме глаза, я бы не обучал физике, химии модель. Но если это какая-то всесторонняя ассистентка, которая из различных областей специалистом будет подсказывать правильное решение, то чем больше обучающая выборка, тем более она всесторонняя, тем результат будет, на мой взгляд, более эффективен.

Сергей Викторович Вишняков:

Но согласитесь, что необязательно учить нейросеть так, как нас учили в школе. Необязательно сначала изучать арифметику, алфавит. Есть вещи, которые, наверное, надо изучать сам процесс.

Павел Романович Варшавский:

Скорее всего, это просто база общих знаний должна быть подключена. Возможно, к ним обратиться. Например, в динамике надо смотреть, как изменяется сетчатка у одного пациента после годичного прихода и снятия этих же показаний. Сразу появится динамический процесс, который, может быть, уже выглядит не просто как

классификация, а мы видим прогресс. У него деградируется сетчатка или, наоборот, все стабильно. То есть здесь могут потребоваться какие-нибудь элементы работы со временем. Здесь скорее не то, что мы его как школьника будем учить, как ребенка, хотя, наверное, такой вариант возможен. Скорее, то, что сказал Сергей Викторович, это возможность у этих помощников как-то выходить за грани только свои узкоспециализированные задачи. Зачем? В случае, если меняется или расширяется область тех задач, которые перед ним стоят. Иначе это просто надо переобучать или переделывать ассистента. Если мы хотим какого-то ассистента, который может не только глаза лечить, а еще и может снимки. А с ИИ так и будет происходить. Как только мы хорошо научились одну задачу решать. Вот сейчас знаки автомобильные распознаются на довольно высокой точности. Начали динамику смотреть, теперь смотрят как там аварии и так далее.

Олег Вадимович Глухов:

При этом есть проблема, когда, допустим, какой-нибудь мужик одевает футболку с такими знаками...Машина не понимает, что это мужик, а не столб.

Павел Романович Варшавский:

Об этом и речь, что надо смотреть не только про нейросети, а вообще про инструменты ИИ, которые можно использовать. Скорее всего, конечно, будет соблазн у человека все время делегировать или сделать его функциональные возможности все шире и шире. И тогда, наверное, вот именно в этом аспекте. А для решения конкретных задач здесь, наверное, не стоит. Сергей Викторович прав, можно ограничиться, если мы заранее знаем, что мы хотим от него. Но просто, как с тем примером, который нам привели, что вот есть человек вырос, да, и вот он решает задачи в этой предметной области, и он до этого изучал все вот это, да, можно ли считать, что нейросеть надо научить всему тому же, да, и вот она тогда будет наравне.

Роман Сергеевич Куликов:

Не наравне, нет. Мы все-таки изначально, я говорю, что целеполагание и оценка результатов – это человеческая функция. Мы говорим про некоторую автоматизацию что можно автоматизировать.

Павел Романович Варшавский:

Здесь скорее-то вот будет зависеть от того, как мы это видим. Что должно быть заложено в систему, или какие возможности мы должны предоставить, чтобы она могла эффективно решать вот ту задачу, которую мы ставим.

Роман Сергеевич Куликов:

Ну, например, эксплуатация известного, или там, сильно известного объекта, в значительной мере известного объекта, технического, например, или какого-то еще, она, наверное, действительно, может быть, не требует какой-то общего, гуманитарной, общечеловеческой подготовки, потому что объект ограничен, значит,

его как-то вписали там, при необходимости дополнили. Но когда идет речь, например, об ассистировании, проектировании того, чего еще не было, с целой кучей сложных ограничений, где, например, выбор решения, он вообще лежит, необязательно в какой-то одной области он лежит, так что вывод, что на это именно вводить нельзя делать, надо переходить на другие типы вещей, и там, заранее это не предусмотрено все случаи жизни, то есть, может быть, в этом-то случае и нужно.

Сергей Викторович Вишняков:

Нет, может, я Ваш вопрос неправильно понял, потому что я думал, что речь идет именно о дублировании пути обучения человека, то есть, должны ли мы читать сказку «Колобок», нейросети сначала? (отвечают: Ну, это следующий вопрос). Нет, для человека это обязательный процесс, потому что не прочитав сказки в детстве, мы не получаем определенных культурных навыков, и мы не можем учиться дальше, для нас это обязательно.

Олег Вадимович Глухов:

Не, ну, может, такая этическая проблема тоже встанет. Александр Павлович, а Вы свое мнение выскажите по поводу этого вопроса?

Александр Павлович Еремеев:

Да, вот я бы хотел сказать, вот в контексте именно... В процессе обучения вот мне кажется, что надо всем студентам примерно в третьего курса читать курсы основы искусственного инженера. Например, Дмитрий Александрович Поспелов четко подтвердил, что искусственный интеллект — это мультидисциплина. То есть и кибернетика, и системный анализ, и психология, и нейрофизиология, вот они, которые о человеческом мышлении, которого мы еще не очень хорошо понимаем сами, но известные психологические факторы.... То есть человек может одновременно оценивать ..., параметры, образы и так далее. Больше сложно, да, есть люди, которые могут запоминать число Π , но это скорее всего патология, чем норма. И вот обучая основам искусственного интеллекта, надо, чтобы люди четко понимали, что это мультидисциплина. И строить этот курс надо по-другому направлению. Скажем, каким-то факультетом, специальностью достаточно подавать искусственный интеллект как средство... Там решение несложных задач, скажем, сокращение поискового пространства. Здесь люди должны понимать, во-первых, что такое рассуждение. Что такое достоверное рассуждение и правдоподобное, индуктивное, адаптивное и так далее. Здесь надо учить студентов пользоваться методами искусственного интеллекта. И вторая группа направлений, это вот касается нашего факультета. Касается факультета... Касается людей, которые могут разрабатывать методы искусственного интеллекта. И разрабатывать, и рекомендовать их для решения прикладных задач в своей предметной области. Это более сложная система. Люди должны понимать, что такое знание, что такое интерпретация, что такое база знаний. Как проверять на корректность и так далее. К сожалению, сейчас наблюдаются ситуации, когда студенты бездумно используют просто нейронные

сети. Но это естественно Как робота обучили, так он и отвечает. Если мы начнем задавать вопросы, связанные там, ну, скажем, со словом СВО и так далее, то, естественно, обученная нейронная сеть на западных выборках, естественно, будет давать ту интерпретацию ответа, совершенно непохожую на нашу. Вот здесь тоже нужно объяснить людям, что искусственный интеллект, это не панацея от всего. Ну, опять-таки, вот когда мы узнаем основу искусственного интеллекта, вот эти аспекты надо будет осветить. А далее уже ориентировать студентов в более старших курсах или на использование метода искусственного интеллекта для решения своих задач, или не только использовать, но и разработать соответствующие методы и системы для своих приложений... И здесь нужно подключать специальные дисциплины, включая и базу дисциплин, математическая, логика, программирование, системная ... и так далее. Вот этого хотелось бы, чтобы тоже было учтено, когда мы будем говорить уже о курсах искусственного интеллекта такого широкого применения.

Олег Вадимович Глухов:
Александр Павлович, спасибо!

Людмила Георгиевна Голубкова:

Есть одно соображение. Я уже говорила в прошлый раз, что я, наверное, среди вас единственный нетехнический специалист. Я в анамнезе лингвист и методист, специалист по интенсивному обучению взрослых. И если мы говорим об обучении студентов и построении программ, это одна задача, и она более-менее понятна. Об этом очень хорошо сказал сейчас наш коллега. А вот педагогика самого искусственного интеллекта, она как раз является дискуссионной и является таковой уже лет как минимум 40. Потому что если вы вспомните Марвина Минского, то он посвятил вот такую книгу тому, чтобы проследить, как ребёнок маленький обучается строить пирамидки. И такие простые из кубиков. Ну, вы читали наверняка. Из кубиков разной конструкции. И в конце книги Минский, он тогда был заведующим лабораторией Массачусетского технологического института, как раз который изучали этот искусственный интеллект. Он приходит к выводу, что машину так как ребёнка обучить нельзя. Но любому человеку, который закончил вуз педагогический, это понятно априори. Что машина — это машина, а человек — это человек. И не потому, что у машины или у человека нет каких-то качеств, которые, соответственно, есть у другой сущности. А просто потому, что, как думает человек, модели именно мышления человека до сих пор нет. Вот изучили мозг уже до, по-моему, нанометров. Как он устроен, что в какой момент включается. А вот как мы думаем, модели нет. Современная наука, в смысле science, это вот наука, которая появилась где-то лет 500 назад, она отрезала определённую часть методов, которые позволяют понять, как человек думает. И пока этого не будет, а в этой парадигме этого не будет, всё время будет аппроксимация, то есть подход к снаряду всё ближе и ближе, но имитировать мышление человека просто невозможно, потому что нет модели. Кто-то со мной может не согласиться, но вот если после семинара вы подумаете, как и о чём вы думали, вы не сможете построить эту модель в том языке, в той системе

представления, который нам через обучение и образование дают. Если закладывать эту модель обучения в машину, то мы можем получить довольно парадоксальный результат, что машина будет тоже как человек накапливать знания, но у машины никогда не будет умений, то есть способности человеческой эти знания превращать в опыт. Вот здесь часто звучало, что у машины есть опыт. Нет у нее опыта. У нее есть таблица и веса, грубо говоря. И довольно такой хитроумный способ, опять же созданный человеком, как это применять к получению машинного результата. К чему я это говорю? Мне кажется, нам стоило бы, ну это было как бы негатив, а сейчас позитив, что подумать о следующей интересной ситуации. Вот учить машину всему вот тому культурному багажу, которому учат в семье, в школе, может быть и не стоит, потому что культурная рамка меняется. И вот те самые студенты, которые приходят, и школьники уже, и они прекрасно могут использовать верхнеуровневые, очень простенькие модели для своих нужд. Коренным образом вот этот культурный, именно культурный сдвиг, он противоречит нашему, замечу, политехническому среднему общему образованию. Только в Германии и у нас было внедрено вот это всеобщее, среднее, замечу, политехническое, не гуманитарное, политехническое образование, на котором вот мы все до сих пор выезжаем, несмотря на все, значит, коллизии предыдущих уже практически 40 лет, когда это рушили. Но все равно это была очень серьезная основа. Вот меняется культурная парадигма. Здесь и сейчас нужен результат, нужно содействие, нужна помощь машины. Я специально говорю слово машина, не нейросеть, чтобы исключить слово мышление, опыт. Вот от этого мы уходим. И если мы ставим такую очень прагматическую задачу, прагматика, это нацеленность на получение результата здесь и сейчас. Я это слово, прагматика, так использую. Вот этот прагматический результат, чтобы получить, значит, машина должна иметь определенную базу знаний. И очень интересно, спасибо, про базу прецедентов. И быть настроена так, чтобы вот решать сначала очень простые задачи, но решать не просто хорошо, а решать предсказуемо, решать так, чтобы количество неправильных, ну или то, что сейчас называется галлюцинации, было минимальным. И чтобы удовлетворять запросы тех людей, которые приходят уже с определенным культурным багажом. Культура – это не театр и книги. Культура – это все, что нас окружает. И DeepSeek – это тоже культурный багаж, как это ни странно звучит. И есть еще одна очень важная задача – «загнать» в эти сети как можно больше материалов по различным областям чувствительным, потому что они исчезают из открытого интернета, из открытого доступа. Они исчезают не просто потому, что не работают сайты, а просто потому, что закрывают. Вот архив.org, он уже практически закрыт. Значит, в США закрыли с 1 июля доступ к национальным архивам. Просто закрыли. Его нельзя получить за деньги. И в той области, в которой вы работаете – энергетика, радиоэлектроника, электроника, радиотехника, математика – все меньше и меньше будет доступ. Это специально делается, чтобы конкурентное преимущество получить. Поэтому все, что можно взять из открытого доступа, должно быть взято. И как можно быстрее, потому что каждый месяц будут закрываться архивы. Журналы останутся, но там будет 100 тысяч китайцев, которые друг у друга переписывают с помощью той же нейросети. И скорее всего, вот эта дельта добавления нового знания будет

минимальной. Сети — вот это нужно, потому что эти специальные знания и будут обуславливать вот тот добавленный результат, про который тоже вы правильно говорили, когда нейросеть на междисциплинарном уровне сможет работать. Если у нее не будет вот этих вот архивов, но она будет рассуждать в терминах учебника физики, Пёрышкин. Тоже хорошая вещь, но для решения профессиональных задач этого будет недостаточно. От тех людей, которые учились по учебнику Пёрышкина, я думаю, уже в этой комнате-то и нет. Потому что я, наверное, самая старшая из вас рискну предположить. У меня уже там были другие авторы, и у меня была другая геометрия, другая алгебра. А вот та, по которой сейчас учатся все американцы, потому что они-то Пёрышкина перевели. И Киселева учебник математики они перевели и в хороших школах используют. А у нас этого уже нет. Уже нет больше, скоро будет уже 60 лет.

Олег Вадимович Глухов:

Спасибо большое. И вот последнее, о чем хотелось поговорить, это выбор тематики следующего семинара. Есть ли у кого-то желание именно какую-нибудь тему предложить?

Константин Александрович Петров:

На самом деле у меня вот, я сразу пришел, сказал, что есть предложение от Сергея Ивановича провести семинар у нас (в НИИСИ), выездной.

Роман Сергеевич Куликов:

Тематики есть и докладчики?

Константин Александрович Петров:

Тематики, докладчики есть.

Роман Сергеевич Куликов:

Какие?

Константин Александрович Петров:

То, что вкратце я говорил на словах в прошлый раз, это применение искусственного интеллекта в области проектирования микроэлектроники.

Роман Сергеевич Куликов:

Какие есть результаты? Текущие? Или задумки?

Константин Александрович Петров:

Нет, нет, нет, есть текущие результаты и задумки. Текущие результаты это, ну просто, грубо говоря, есть этап-маршрут проектирования. Да, есть определенные моменты, которые посчитали нужным закрыть, которые практически применимы. Но, честно говоря, я вот сейчас в конце этого семинара начинаю немножко сомневаться, потому

что мы настолько глубоко в это все лезем. Сейчас мы рассмотрели три из четырех тематик, да, даже не доклад. Это один доклад, и он четыре тематики. Рассмотрели три. Если мы будем, допустим, три-четыре доклада представлять, если мы будем на таком уже уровне обсуждения, который у нас сейчас здесь, то больше одного доклада, мне кажется, нормально будет.

Роман Сергеевич Куликов:

Задача семинара, это не галочка получить.

Константин Александрович Петров:

Нет, нет, что нам интересно в этом семинаре, то, что мы очень глубоко с вами копаем.

Роман Сергеевич Куликов:

Так он для того и сделан, чтобы разобраться.

Олег Вадимович Глухов:

Так, это предлагается одна большая тематика в рамках семинара?

Константин Александрович Петров:

В целом направление, то есть по раз, два, три доклада. Причем, часть из них — это открытые доклады, часть из них надо будет поговорить с заказчиком.

Олег Вадимович Глухов:

Хорошо.

Роман Сергеевич Куликов:

Как минимум, какую-то тему под семинар из области проектирования микросхем.

Константин Александрович Петров:

Да, но я просто думаю, надо будет обсудить, может даже вам с Сергеем Ивановичем лучше, потому что он, конечно, все смотрит очень позитивно.

Роман Сергеевич Куликов:

Как направление семинара принимается, наверное, Олег, как ты считаешь?

Олег Вадимович Глухов:

Давайте еще обсудим, потому что есть идея от компании Инноцифра. Они там буквально хотят про свой опыт очень подробно рассказать. Это компания, которая внедряет в московские техникумы гибкие образовательные траектории для студентов.

Роман Сергеевич Куликов:

Педагогика в том смысле, что они ассистируют педагогу.

Константин Александрович Петров:
У нас тоже такое есть.

Павел Романович Варшавский:
Вопрос, который сейчас подняли, надо четко понимать. Педагогика ИИ, за что мы беремся? Я полностью согласен это просто две диаметрально разные задачи, и тут даже попыток учить как бы машину как человека — это бесполезное дело. Не то, что там известно, скорее принципиально разные подходы.

Роман Сергеевич Куликов:
Я интересуюсь тем, чтобы учить машину, чтобы она повышала продуктивность разработчиков. Это мой интерес в семинаре. То, что Олег говорит, там коллеги другой задачей занимаются, они ассистируют педагогу.

Олег Вадимович Глухов:
Мое мнение, что направление использования ИИ в образовании мы будем на семинаре касаться часто, так как сами в ВУЗе находимся, а также всегда будет подниматься вопрос о том, а как учить современного студента, в том числе и тому, как пользоваться разрабатываемыми ИИ-ассистентами. У коллег из Инноцифры довольно интересный и успешный кейс, поэтому не вижу здесь противоречий. Давайте рассмотрим их опыт на ближайших семинарах.
Коллеги, спасибо всем за участие, завершаем!