



Ассоциация
Российских
Банков



Национальный исследовательский
институт Доверия, Достоинства и Права

КОГНИТИВНАЯ БЕЗОПАСНОСТЬ

Материалы «Рабочего завтрака у Тосуняна»
12 июля 2025 года

ДОКЛАДЧИКИ:



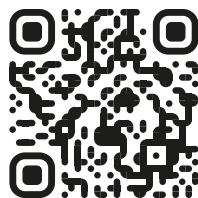
Кефели Игорь Федорович

доктор философских наук, профессор,
ведущий научный сотрудник
Северо-Западного института управления – филиала
РАНХиГС при Президенте РФ,
профессор Высшей школы международных отношений
Санкт-Петербургского университета Петра Великого,
заместитель главного редактора журнала
«Евразийская интеграция: экономика, право, политика»,
эксперт РАН



Колбанёв Михаил Олегович

доктор технических наук, профессор,
профессор кафедры информационных систем
и технологий Санкт-Петербургского
государственного экономического университета



НКС ООН РАН
Научно-консультативный совет
по правовым, психологическим
и социально-экономическим проблемам общества
Отделения общественных наук РАН

АРБ
Ассоциация российских банков

НИИ ДДнП
Национальный исследовательский институт
Доверия, Достоинства и Права

Когнитивная безопасность

Материалы заседания 12 июля 2025 года

Под общей редакцией
академика РАН
Г.А. Тосуняна

Москва
2026

УДК [004.056+004.8]:316.6](063)
ББК 32.813я431+32.973.26-018.2я431
К57

Когнитивная безопасность : материалы заседания 12 июля 2025 года / Научно-консультативный совет по правовым, психологическим и социально-экономическим проблемам общества Отделения общественных наук Российской академии наук ; Ассоциация российских банков ; Национальный исследовательский институт Доверия, Достоинства и Права ; [под общ. ред. академика РАН Г.А. Тосуняна]. – М.: ООО «Новые печатные технологии», 2026. – 144 с. – ISBN 978-5-6056500-3-4

В условиях наступившей цифровой эпохи, стремительного развития искусственного интеллекта, когда когнитивные способности человека перегружены потоком избыточной информации, возрастают риски, которые требуют внимания.

В сборнике обозначен круг проблем когнитивной безопасности общества и предложены некоторые меры, направленные на противодействие угрозам, включая соблюдение этических стандартов в области искусственного интеллекта, разработку законов и нормативов для его защищенного использования и другие.

Читатель узнает о возможных последствиях хайпа, который складывается вокруг темы искусственного интеллекта, а также о набирающих силу когнитивных войнах и их негативном воздействии.

УДК [004.056+004.8]:316.6](063)
ББК 32.813я431+32.973.26-018.2я431

Охраняется в соответствии с международным правом и российским законодательством об авторском праве.

ISBN 978-5-6056500-3-4

© Тосунян Г.А., составление, 2026

СОДЕРЖАНИЕ

Состав Научно-консультативного совета по правовым, психологическим и социально- экономическим проблемам общества (НКС ППСЭПО) ООН РАН	5
Справка	8

ВСТУПИТЕЛЬНОЕ СЛОВО

акад. ТОСУНЯН Г.А.	11
---------------------------	-----------

Доклад 1 д. филос. н., проф. КЕФЕЛИ И.Ф.	20
---	-----------

КОГНИТИВНАЯ БЕЗОПАСНОСТЬ КАК СИСТЕМА: ПРАВО НА СУЩЕСТВОВАНИЕ

к. филос. н. КРЖЕВОВ В.С. – д. филос. н. КЕФЕЛИ И.Ф.	33
д. филос. н. СОРИНА Г.В. – д. филос. н. КЕФЕЛИ И.Ф.	35
д. э. н. КРИВОВ В.Д. – д. филос. н. КЕФЕЛИ И.Ф.	39
д. п. н. СЕРГЕЕВ С.Ф. – д. филос. н. КЕФЕЛИ И.Ф.	41

Доклад 2 д. т. н., проф. КОЛБАНЁВ М.О.	48
---	-----------

ИСКУССТВЕННЫЙ ИНТЕЛЛЕКТ КАК ХАЙП

к. филос. н. ЕФРЕМОВ О.А. – д. т. н. КОЛБАНЁВ М.О.	71
к. б. н. ХАЛЯВКИН А.В. – д. т. н. КОЛБАНЁВ М.О.	73
д. т. н. КОНЯВСКИЙ В.А. – д. т. н. КОЛБАНЁВ М.О.	76
СМОЛИН В.С. – д. т. н. КОЛБАНЁВ М.О.	77
д. э. н. КРИВОВ В.Д. – д. т. н. КОЛБАНЁВ М.О.	80
к. филос. н. ЗАЙКОВА А.С. – д. т. н. КОЛБАНЁВ М.О.	82
акад. БАРАНОВСКИЙ В.Г. – д. т. н. КОЛБАНЁВ М.О.	86

МОСКОВКИН Л.И.	88
д. т. н. КОНЯВСКИЙ В.А. – д. п. н. СЕРГЕЕВ С.Ф.	89
к. филос. н. ЕФРЕМОВ О.А. – акад. ТОСУНЯН Г.А.	98
д. э. н. КРИВОВ В.Д.	103
д. филос. н., проф. СОРИНА Г.В.	105
д. б. н. КАВТАРАДЗЕ Д.Н.	107
СМОЛИН В.С.	110
д. п. н. СЕРГЕЕВ С.Ф.	113
к. филос. н. КРЖЕВОВ В.С. – акад. ТОСУНЯН Г.А.	115
к. б. н. ХАЛЯВКИН А.В.	118
МОСКОВКИН Л.И.	120
 ЗАКЛЮЧИТЕЛЬНЫЕ СЛОВА	
д. филос. н., проф. КЕФЕЛИ И.Ф.	121
д. т. н., проф. КОЛБАНЁВ М.О.	123
 ЗАКЛЮЧЕНИЕ	
акад. ГУСЕЙНОВ А.А.	125
акад. ТОСУНЯН Г.А.	131
 Список литературы, опубликованной по итогам заседаний НКС ООН РАН, НИИ ДДиП и «Открытых дискуссий» президента АРБ	 136

**СОСТАВ НАУЧНО-КОНСУЛЬТАТИВНОГО СОВЕТА
ПО ПРАВОВЫМ, ПСИХОЛОГИЧЕСКИМ И СОЦИАЛЬНО-
ЭКОНОМИЧЕСКИМ ПРОБЛЕМАМ ОБЩЕСТВА
(НКС ППСЭПО) ООН РАН**

СОПРЕДСЕДАТЕЛИ:

ГУСЕЙНОВ
АБДУСАЛАМ
АБДУЛКЕРИМОВИЧ

академик, д. филос. н., научный руководитель Института философии РАН

КОКОШИН
АНДРЕЙ
АФАНАСЬЕВИЧ

академик, д. и. н., директор Центра перспективных исследований национальной безопасности России
Экспертного института НИУ ВШЭ

ТОСУНЯН
ГАРЕГИН
АШОТОВИЧ

академик, д. ю. н., президент Ассоциации российских банков

УЧЕНЫЙ СЕКРЕТАРЬ:

РЕДЬКО
НИКОЛАЙ
ВИТАЛЬЕВИЧ

к. э. н., эксперт Национального исследовательского института Доверия, Достоинства и Права

ЧЛЕНЫ НАУЧНОГО СОВЕТА:

АВETИСЯН
АРУТЮН
ИШХАНОВИЧ

академик, д. ф.-м. н., директор Института системного программирования им. В.П. Иванникова РАН

АГАНБЕГЯН
АБЕЛ
ГЕЗЕВИЧ

академик, д. э. н., профессор

АПОЛИХИН
ОЛЕГ
ИВАНОВИЧ

чл.-корр., д. м. н., директор НИИ урологии и интервенционной радиологии им. Н.А. Лопаткина (филиал ФГБУ «НМИЦ радиологии» Минздрава России)

АУЗАН
АЛЕКСАНДР
АЛЕКСАНДРОВИЧ

д. э. н., декан экономического факультета МГУ им. М.В. Ломоносова

БАТУРИН
ЮРИЙ
МИХАЙЛОВИЧ

чл.-корр., д. ю. н., главный научный сотрудник отдела методологических и междисциплинарных проблем развития науки Института истории естествознания и техники им. С.И. Вавилова РАН

БУЗНИК
ВЯЧЕСЛАВ
МИХАЙЛОВИЧ

академик, д. х. н., заместитель академика-секретаря ОХНМ РАН, начальник лаборатории Всероссийского НИИ авиационных материалов

ГРАЧЕВА
ЕЛЕНА
ЮРЬЕВНА

д. ю. н., профессор, первый проректор ФГБОУ ВО «Московский государственный юридический университет им. О.Е. Кутафина» (МГЮА)

ГРИНБЕРГ
РУСЛАН
СЕМЕНОВИЧ

чл.-корр., д. э. н., научный руководитель Института экономики РАН

ДАНИЛОВ-ДАНИЛЬЯН
АНТОН
ВИКТОРОВИЧ

к. э. н., заместитель председателя Общероссийской общественной организации «Деловая Россия»

ЕРМАКОВА
ЖАННА
АНАТОЛЬЕВНА

чл.-корр., д. э. н., профессор, заведующий кафедрой банковского дела и страхования ФГБОУ ВО «Оренбургский государственный университет»

ЖУРАВЛЕВ
АНАТОЛИЙ
ЛАКТИОНОВИЧ

академик, д. п. н., научный руководитель Института психологии РАН

ИВАНОВ
ВИЛЕН
НИКОЛАЕВИЧ

чл.-корр., д. филос. н., главный научный сотрудник Института социально-политических исследований ФНИСЦ РАН

ИЛЬИН
ВЛАДИМИР
АЛЕКСАНДРОВИЧ

чл.-корр., д. э. н., профессор, научный руководитель Вологодского научного центра РАН

КАСАВИН
ИЛЬЯ
ТЕОДОРОВИЧ

чл.-корр., д. филос. н., руководитель сектора социальной эпистемологии Института философии РАН

КЛЕПАЧ
АНДРЕЙ
НИКОЛАЕВИЧ

к. э. н., главный экономист ВЭБ.РФ

ЛЕКТОРСКИЙ
ВЛАДИСЛАВ
АЛЕКСАНДРОВИЧ

академик, д. филос. н., главный научный сотрудник Института философии РАН

МЕДВЕДЕВ
ПАВЕЛ
АЛЕКСЕЕВИЧ

д. э. н., профессор

МИРКИН
ЯКОВ
МОИСЕЕВИЧ

д. э. н., профессор

НЕСТИК
ТИМОФЕЙ
АЛЕКСАНДРОВИЧ

д. п. н., профессор РАН, заведующий лабораторией социальной и экономической психологии Института психологии РАН

НИГМАТУЛИН
РОБЕРТ
ИСКАНДРОВИЧ

академик, д. ф.-м. н., научный руководитель Института океанологии им. П.П. Ширшова РАН

ПЕТРЕНКО
ВИКТОР
ФЕДОРОВИЧ

чл.-корр., д. п. н., заведующий лабораторией психологии общения факультета психологии МГУ им. М.В. Ломоносова

ПОГОСЯН
ГЕВОРК
АРАМОВИЧ

академик Национальной академии наук Армении (НАН РА), иностранный член РАН, д. социол. н., научный руководитель Института философии, социологии и права НАН РА

САВЕНКОВ
АЛЕКСАНДР
НИКОЛАЕВИЧ

академик, д. ю. н., директор Института государства и права РАН

САННИКОВА
ЛАРИСА
ВЛАДИМИРОВНА

д. ю. н., профессор РАН, руководитель Центра правовых исследований цифровых технологий Государственного академического университета гуманитарных наук

САРКИСЯН
ТИГРАН
СУРЕНОВИЧ

к. э. н., заместитель председателя правления
Евразийского банка развития

СМИРНОВ
АНДРЕЙ
ВАДИМОВИЧ

академик, д. филос. н., главный научный сотрудник
Института философии РАН

СОЛОДКОВ
ВАСИЛИЙ
МИХАЙЛОВИЧ

к. э. н., директор Банковского института НИУ
ВШЭ

ТЕДЕЕВ
АСТАМУР
АНАТОЛЬЕВИЧ

д. ю. н., профессор кафедры государственного
аудита Высшей школы государственного аудита
(факультет) МГУ им. М.В. Ломоносова

ТОРШИН
АЛЕКСАНДР
ПОРФИРЬЕВИЧ

к. ю. н., действительный государственный советник
РФ I класса

ТОЩЕНКО
ЖАН
ТЕРЕНТЬЕВИЧ

чл.-корр., д. филос. н., профессор, главный научный
сотрудник Института социологии ФНИСЦ
РАН

УГРЮМОВ
МИХАИЛ
ВЕНИАМИНОВИЧ

академик, д. б. н., заведующий лабораторией нервных
и нейроэндокринных регуляций Института
биологического развития им. Н.К. Кольцова РАН

УШАКОВ
ДМИТРИЙ
ВИКТОРОВИЧ

академик, д. п. н., директор Института психологии
РАН

ХАБРИЕВА
ТАЛИЯ
ЯРУЛЛОВНА

академик, д. ю. н., директор Института законодательства
и сравнительного правоведения при Правительстве России

ЧЕРЕШНЕВ
ВАЛЕРИЙ
АЛЕКСАНДРОВИЧ

академик, д. м. н., научный руководитель Института
иммунологии и физиологии Уральского отделения РАН

ЧЕРНЫШ
МИХАИЛ
ФЕДОРОВИЧ

чл.-корр., д. социол. н., научный руководитель Федерального
научно-исследовательского социологического центра РАН

ЧЕХОНИН
ВЛАДИМИР
ПАВЛОВИЧ

академик, д. м. н., вице-президент РАН, заведующий
кафедрой медицинских нанотехнологий медико-биологического
факультета Российского государственного медицинского
университета им. Н.И. Пирогова

ШАБУНОВА
АЛЕКСАНДРА
АНАТОЛЬЕВНА

д. э. н., директор Вологодского научного центра
РАН

ЭКМАЛЯН
АШОТ
МАМИКОНОВИЧ

д. филос. н., профессор

ЮРЕВИЧ
АНДРЕЙ
ВЛАДИСЛАВОВИЧ

чл.-корр., д. п. н., заместитель директора по научной
работе Института психологии РАН

СПРАВКА

- о НКС ООН РАН (Научно-консультативном совете по правовым, психологическим и социально-экономическим проблемам общества Отделения общественных наук),
- о НИИ ДДиП (Национальном исследовательском институте Доверия, Достоинства и Права),
- о «Рабочем завтраке у Тосуняна»,
- о проекте «Открытые дискуссии президента АРБ»
и об этом издании

1. НКС ООН РАН был создан в 2012 году как Совет по правовым, экономическим, социально-политическим и психологическим аспектам финансово-кредитной системы.

Заседания Совета проводились в Отделении общественных наук РАН два раза в год.

В феврале 2020 года члены НКС приняли решение расширить компетенцию Совета, перейдя от рассмотрения вопросов развития финансового рынка к более широкому кругу проблем развития общества, поставив во главу угла своих исследований и дискуссий вопросы: «В каком обществе мы живем? Какое общество мы хотели бы оставить своим потомкам в наследство?»

И в сентябре 2021 года постановлением Президиума РАН Совет был преобразован в Научно-консультативный совет по правовым, психологическим и социально-экономическим проблемам общества ООН РАН.

Сопредседателями Совета стали академики РАН А.А. Гусейнов, А.А. Кокошин и Г.А. Тосунян.

2. С середины 90-х годов по субботам раз в две-три недели в Ассоциации российских банков проходят «Рабочие завтраки у Тосуняна», в которых принимали и принимают участие банкиры, представители ЦБ, Госдумы, Совета Федерации, различных ведомств, академической науки, вузов, эксперты по финансово-банковскому профилю.

Каждый «Рабочий завтрак у Тосуняна» (далее – «Рабочий завтрак») проходит по заранее согласованной повестке дня и с заявленными докладчиками.

На них до недавнего времени обсуждались преимущественно проблемы экономики, финансовой сферы, нормативно-правовые акты, регулирующие эту сферу. Но в ряде случаев и другие вопросы развития общества.

В последние годы спектр вопросов, рассматриваемых на «Рабочих завтраках», и круг экспертов заметно расширились.

Этому во многом способствовало участие в них известных ученых.

Характерной особенностью «Рабочих завтраков» было и остается то, что они проходят с завидной регулярностью по субботам в 9.00 утра и зимой, и летом, и даже 31 декабря. Их продолжительность примерно 3–4 часа.

3. В конце 2019 года был учрежден Национальный исследовательский институт Доверия, Достоинства и Права (НИИ ДДиП).

Это частный институт, целью которого, если вкратце, является многогранное изучение вопросов человеческой жизнедеятельности и общественных процессов, которые наибольшим образом влияют на развитие доверия в обществе, повышение ответственности и чувства собственного достоинства у граждан страны и на формирование уважения друг к другу.

Институт приступил к работе в начале 2020 года в формате научных заседаний с коллегами, интересующимися проблемами доверия и достоинства, их правового обеспечения и стимулирования.

Иначе говоря, институт пригласил на общественных началах работать на его площадке всех, кто желает внести свою лепту в изменение траектории движения общества «войны всех против всех» в сторону общества «доверия, достоинства и уважения друг к другу»!

4. В конце марта 2020 года был объявлен локдаун.

Встал вопрос: заморозить на какое-то время работу НКС ООН, НИИ ДДиП, АРБ и «Рабочие завтраки у Тосуняна»?

Или искать какое-то другое решение?

Тогда же возникла идея, что заседания НКС ООН, НИИ ДДиП и «Рабочие завтраки» можно объединить, используя онлайн-формат.

Проанализировав практику последних лет, мы с коллегами пришли к выводу, что довольно часто и на заседаниях НКС, и на «Рабочих

завтраках», и на заседаниях Института мы поднимаем и обсуждаем схожие вопросы.

Было принято решение начать проводить совместные заседания.

За прошедшее с апреля 2020 года время было проведено 152 «Рабочих завтрака у Тосуняна», большинство из которых прошло в очно-заочной форме.

Примерно 20 человек лично присутствовали на завтраках, а остальные, от 50 до 100 и более человек, принимали участие в режиме Zoom, видя, слыша «живых» участников и докладчиков, также присоединялись к дискуссии.

В последующем по видеозаписи каждое заседание стенографировалось с тем, чтобы можно было издать материалы этих дискуссий.

В настоящее время накопился огромный объем материалов для публикаций, и мы начали их издание в виде представленных вашему вниманию сборников.

5. С 2013 года Ассоциация российских банков ведет проект «Открытые дискуссии президента АРБ».

Проект направлен на обсуждение широкого круга экономических, правовых, философских, социально-психологических и других актуальных проблем развития нашего общества и на развитие культуры дискуссии в целом. Спикерами «Открытых дискуссий президента АРБ» (далее – «Открытые дискуссии») выступают известные ученые, общественные деятели и представители бизнеса.

Вузами-партнерами проекта являются более 90 российских вузов, расположенных на территории всей России – от Владивостока до Калининграда.

Как правило, в каждой «Открытой дискуссии» дистанционно участвуют от 40 до 90 вузов. Численность интернет-аудитории в среднем составляет около 2 тыс. человек.

Последние два года «Открытые дискуссии» проводятся ежемесячно.

За 10 лет состоялось 95 дискуссий.

С информацией о прошедших дискуссиях, презентационными материалами спикеров и видеозаписями можно ознакомиться на сайте arb.ru в разделе «Открытые дискуссии».

Г.А. ТОСУНЯН, академик РАН,
президент Ассоциации российских банков

ВСТУПИТЕЛЬНОЕ СЛОВО

ТОСУНЯН Г.А.

акад. РАН

Приветствую всех на совместном заседании Научно-консультативного совета Отделения общественных наук РАН по правовым, психологическим и социально-экономическим проблемам общества и Института Доверия, Достоинства и Права, которое проходит при поддержке Ассоциации российских банков.

Для участия зарегистрировалось 123 человека, из них 33 – представители науки и вузов, в том числе семь членов Академии наук, а также банкиры, государственные и общественные деятели, эксперты.

Несмотря на то, что середина лета, тем не менее мы с Абдусаламом Абдулкеримовичем проводим наши заседания.

Мы никого не принуждаем и не обязываем.

120 с лишним человек даже в летнее время желают участвовать в рабочем завтраке!

Докладчики, находясь в отпуске, заявили свои доклады. Как можно отказать себе и всем нам в такой роскоши общения даже в пик сезона отпусков.

Тема сегодняшнего заседания – «Когнитивная безопасность» – особенно актуальна в нашу цифровую эпоху, поскольку когнитивные способности человека сегодня перегружены потоком избыточной информации.

Исследователи заявляют, что это может вызывать психическое истощение людей. Думаю, не только может, а скорее всего, и вызывает у многих такое состояние.

В конце прошлого века появился термин «синдром информационной усталости». Это дисфункциональное состояние, вызванное чрезмерным потоком информации и перегрузкой системы восприятия человека.

Интересное сопоставление сделал Южно-Калифорнийский университет. Там подсчитали, что объем информации, потребляемой человеком ежедневно, в 1986 году был сопоставим с 40 газетами, а в 2006 году – уже со 174 газетами.

Можно предположить, насколько этот объем вырос еще через 20 лет, то есть к настоящему моменту.

Однако вряд ли целесообразно продолжать проводить эти расчеты в газетах, потому что меняется не только количество, но и формат и особенности восприятия все возрастающих объемов данных.

Цифровизация привела к тому, что большинство взаимодействий у людей происходит через экраны мобильных телефонов и других устройств.

Очевидно, что большие информационные потоки могут как обогащать нас, так и вредить нашему сознанию. И сложно сказать, чего здесь больше – пользы или вреда. Поэтому и возникает тема когнитивной безопасности, и она связана в том числе с процессами цифровизации.

Мы на наших заседаниях рассматривали цифровизацию преимущественно в экономике.

Но цифровизация затрагивает все сферы жизни, даже сферу межличностного общения. А может быть, не даже, а в большей степени сферу межличностного общения, которая является частью и бизнеса, и жизни, и всех форм взаимодействия.

За последние десятилетия формат межличностного общения существенно изменился. Каждый сегодня использует мессенджеры в телефонах, в рабочих процессах все чаще используются видеоконференции, Zoom.

Однако никакой Zoom и никакой WhatsApp, по моему мнению, не заменят непосредственное общение людей глаза в глаза.

С моей точки зрения, ценность личного общения сегодня даже возрастает и ее трудно переоценить.

В прошлый раз наши докладчики-психологи затрагивали эту тему. Если бы мобильный телефон был у одного человека, то он не имел бы смысла. Но, когда он становится всеобъемлющим инструментом общения, тогда он приобретает свою значимость.

Этот переход происходит не сразу. Сначала какая-то часть общества обладает этой возможностью, а потом уже она становится глобальной.

Цифровизация влияет и на мышление человека. Пару-тройку десятилетий назад для того, чтобы получить знания, нужно было проделать большой путь, например, пойти в библиотеку.

Я, будучи аспирантом и после, очень любил библиотеку имени Ленина и практически раза четыре в неделю вечерами ходил туда, чтобы изучать необходимую литературу и найти нужную для диссертации информацию. Скажу вам, это был трудоемкий процесс – заказать и получить нужные книги и материалы.

Сегодня мы можем моментально получить нужные данные на экране смартфона. И это будут в большинстве своем актуальные и качественные данные. Правда, иногда, конечно, могут быть и всякого рода фальшивки и не очень качественная информация, но чаще всего необходимая информация достоверна.

Это, конечно, можно и нужно считать прогрессом. Но есть и обратная сторона такой тотальной цифровизации.

Информацию мы получаем не всегда целенаправленно и осознанно, то есть добровольно.

Наш мозг постоянно атакуют терабайты данных: бесконечные ленты новостей, социальные сети, мессенджеры, реклама в ее многочисленных проявлениях.

Все эти источники создают информационный шум, причем такой сильный шум, который перегружает сознание и может затруднять концентрацию внимания на важных задачах. Это влияет и на эмоции, и на принятие решений, и в конечном счете, конечно, влияет на поведение людей.

Информации так много, что люди часто перестают критически ее осмысливать. И этим, кстати, очень успешно пользуются мошенники.

Для финансовой сферы это большая проблема – телефонное мошенничество и социально-криминальная инженерия. Мы об этом также не раз говорили.

Мошенники хорошо подготовлены и используют психологические методы. Обман строится на манипуляциях сознанием в условиях критических ситуаций и перегрузки системы восприятия.

Технологии так продвинулись, что мошенники используют и возможности генерации голоса, и даже фальсификацию видеоизображения.

Я думаю, что наше информационное общество, несмотря на все его технологические преимущества, сегодня переживает в какой-то степени кризис доверия.

Доверие в цифровую эпоху становится очень дефицитным ресурсом.

Возможностей коммуникации много, а доверия к ним становится все меньше. Это происходит, может быть, из-за того, что слишком много коммуникаций.

Мы уже не снимаем трубки, не доверяя незнакомым звонкам. И даже когда звонки вроде как знакомые, все равно бывают коллизии.

Еще одно из ярких проявлений цифровой трансформации – это стремительное распространение технологий искусственного интеллекта.

По этой теме мы провели ряд мероприятий, выпустили два печатных издания.

Это «Искусственный интеллект»¹ – заседание было в октябре 2023 года, мы рассматривали положительные и отрицательные стороны этой технологии.

«Искусственный интеллект в банковской сфере»² – по этой теме мы провели «Открытую дискуссию президента АРБ» в феврале прошлого года.

Эти издания есть и в печатном, и в электронном виде на сайте нашего Совета, так что можете с ними ознакомиться.

Финансовая отрасль – одна из самых продвинутых в использовании новых технологий, в том числе технологий больших данных и искусственного интеллекта. Они уже встроены во многие бизнес-проекты на уровне повседневного применения. Однако это сопряжено с рядом ограничений и рисков, которые мы рассмотрели на упомянутой «Открытой дискуссии» с участием представителей

¹ [Искусственный интеллект \(14.10.2023\) / под общ. ред. академика РАН Г.А. Тосуняна. М.: ООО «Новые печатные технологии», 2024. 182 с.](#)

² [Искусственный интеллект в банковской сфере \(29.02.2024\) / под общ. ред. академика РАН Г.А. Тосуняна. М.: ООО «Новые печатные технологии», 2024. 121 с.](#)

коммерческих банков, Центрального банка, а также Института системного программирования РАН, Института философии РАН.

Технологии искусственного интеллекта – так ли они важны и полезны сейчас?

Если да, то насколько корректно эти технологии причислять к искусственному интеллекту?

Потому что сегодня любую технологию, хоть сколько-нибудь связанную с компьютерной обработкой данных, модно причислять к искусственному интеллекту и нейросетям. Не просто модно, но и прибыльно, этим понятием спекулируют.

Инвесторы позитивно реагируют на любое упоминание искусственного интеллекта.

Многие компании опасаются, что упустят конкурентные преимущества, если не будут внедрять эти технологии, даже если сегодня конкретная польза для них не очевидна.

В принципе, они правильно делают, потому что упустить момент – значит упустить рынок. Действительно, это серьезный риск, но все равно нужно соблюдать баланс.

В 2024 году стоимость общемирового рынка искусственного интеллекта достигла 300 млрд долларов.

Рынок растет примерно на 20% в год.

Но на самом деле большая часть этих технологий является обычными автоматизированными системами, построенными на традиционных алгоритмах.

Ключевое же отличие искусственного интеллекта – это имитация когнитивных функций человека, включая поиск решений без заранее заданного алгоритма, то есть это имитация творческих способностей человека.

Искусственный интеллект получает развитие и финансирование не только в бизнесе, но и на государственном уровне.

В 2019 году издан указ президента России «О развитии искусственного интеллекта в Российской Федерации», утверждены «Национальная стратегия развития искусственного интеллекта на период до 2030 года» и федеральный проект «Искусственный интеллект» с бюджетным финансированием более 24 млрд рублей.

Возможно, искусственный интеллект – это наша «новая нефть». Для нас сейчас важно найти какую-то замену нефти. Или, как сформулировано в теме нашего второго докладчика, может быть, это всего лишь хайп.

Приведу аналогию из финансовой сферы. Цифровой рубль – тоже проект, вокруг которого много ажиотажа и даже спекуляций.

Банковское сообщество занимает сдержанно-критическую позицию по внедрению цифрового рубля.

У банков много вопросов в части механизма внедрения и использования третьей формы национальной валюты. Есть наличная, есть безналичная, теперь предлагается третья форма – цифровой рубль.

Для небольших и региональных банков остро стоит вопрос финансирования этого проекта. Минимальные инвестиции для работы с цифровым рублем значительно превышают ИТ-бюджеты небольших российских банков.

Некоторые представители крупных банков в свою очередь признаются, что не понимают преимуществ и не видят перспектив цифровой валюты внутри страны.

Свою точку зрения на эту тему мы многократно высказывали публично.

Расскажу буквально в двух словах.

Ассоциация российских банков в своих письмах в адрес Центрального банка представила множество аргументов против форсированного внедрения цифрового рубля. Принятое в феврале этого года решение Банка России о переносе сроков проекта свидетельствует о том, что регулятор, слава Богу, принимает во внимание позицию АРБ и банков.

Надо отметить, что мы вовсе не отрицаем важности инноваций на финансовом рынке. Наоборот, мы поддерживаем качество новых услуг для удобства граждан и бизнеса.

Однако мы считаем, что внедрение цифрового рубля должно быть обосновано, просчитано и должно реализовываться без какой-то излишней суеты. Частичное замещение цифровым рублем безналичных денег может существенным образом повлиять на структуру балансов кредитных организаций.

Я не хотел бы сейчас подробно на этом останавливаться. Наша позиция изложена в Годовом докладе Ассоциации российских банков, он есть на сайте Ассоциации в разделе «Съезды АРБ»³.

В части искусственного интеллекта предлагаю действовать так же последовательно и аргументированно.

Во-первых, аккуратно подходить к терминологии, чтобы содержательно обсуждать одно и то же явление, а не разные сущности, потому что мы подменяем иногда понятия и перескакиваем на другие темы.

Во-вторых, прогнозировать последствия повседневного упоминания, или того хайпа, который складывается вокруг темы искусственного интеллекта.

³ [Доклад к съезду Ассоциации российских банков – 2025.](#)

На этом я завершаю свою вводную часть. Теперь перейдем к повестке дня.

Оба докладчика сегодня из Санкт-Петербурга.

Первый вопрос – «Когнитивная безопасность как система: право на существование».

Докладчик – **Кефели Игорь Федорович**,

– доктор философских наук, профессор,

– ведущий научный сотрудник Северо-Западного института управления – филиала РАНХиГС при Президенте РФ,

– профессор Высшей школы международных отношений Санкт-Петербургского университета Петра Великого,

– заместитель главного редактора журнала «Евразийская интеграция: экономика, право, политика».

Второй вопрос – «Искусственный интеллект как хайп».

Докладчик – **Колбанёв Михаил Олегович**,

– доктор технических наук,

– профессор кафедры информационных систем и технологий Санкт-Петербургского государственного экономического университета,

– профессор кафедры информационных систем Санкт-Петербургского государственного электротехнического университета им. В.И. Ленина.

Если наши докладчики с нами, можем начинать. Игорь Федорович, Вам слово.

ДОКЛАД 1

КЕФЕЛИ И.Ф.

д. филос. н., проф., ведущий научный сотрудник Северо-Западного института управления – филиала РАНХиГС при Президенте РФ, профессор Высшей школы международных отношений Санкт-Петербургского университета Петра Великого, заместитель главного редактора журнала «Евразийская интеграция: экономика, право, политика», эксперт РАН

КОГНИТИВНАЯ БЕЗОПАСНОСТЬ КАК СИСТЕМА: ПРАВО НА СУЩЕСТВОВАНИЕ

Спасибо за предоставленную возможность выступить. Для меня большая честь – принять участие в вашем высоком собрании.

К выступлению я пригласил своего товарища, коллегу, соавтора – Михаила Олеговича Колбанёва.

Наша задача – рассказать о направлении развития современной науки, или, точнее, междисциплинарной области под названием когнитология, или когнитивные науки, а также о проблемах развития когнитивных технологий, которые, по сути, должны обеспечивать все направления поддержания когнитивной безопасности – личности, социума, государства.

Этой темой я стал активно заниматься последние пять лет, даже чуть больше. Поначалу для себя, а затем мы провели ряд теоретических семинаров.

И в конечном счете это завершилось изданием серии работ, о которых я скажу несколько слов.

Они связаны с развитием того направления глобалистики, которое касается проблемы глобальной безопасности.

В рамках этой комплексной проблемы я сосредоточил свое внимание и привлек коллектив авторов для того, чтобы выяснить, какое место занимает в системе

глобальной безопасности безопасность информационная, информационно-психологическая и когнитивная.

И в 2017 году совместно с нашим институтом (Северо-Западным институтом управления – филиалом РАН-ХиГС) и Институтом информатики и автоматизации Российской академии наук, вместе с бывшим тогда его директором, членом-корреспондентом РАН Рафаэлем Мидхатовичем Юсуповым, недавно, к сожалению, скончавшимся, мы издали первую монографию «Информационно-психологическая и когнитивная безопасность».

Поначалу был даже некоторый конфуз.

Какая когнитивная безопасность?

Информационно-психологическая есть – ее достаточно.

Но после некоторых дискуссий мы сошлись на том, что все-таки когнитивные операции и когнитивная безопасность – это реально существующие и вполне серьезные вещи.

И мы издали эту работу.

В ней был рассмотрен ряд проблем в постановочном формате: а есть ли когнитивная безопасность и в чем она состоит, в чем ее содержание?

Так коллективная монография «Информационно-психологическая и когнитивная безопасность» стала, по сути дела, первой попыткой разобраться в существе проблемы. Тем более, как выяснилось позже, в те же годы военные специалисты НАТО полным ходом разрабатывали идеологию и технологии ведения когнитивной войны.

Спустя несколько лет, уже в 2023 году, мы провели небольшой теоретический семинар на базе нашего Северо-Западного института управления.

Тогда же созрела идея подготовить коллективную монографию уже с бóльшим количеством авторов.

Один автор, на мой взгляд, может лишь на научно-популярном уровне писать книгу о междисциплинарной области когнитивных наук.

Потому что каждый раздел когнитологии достаточно специфичен, и говорить о нем на уровне любознательности и просветительства было бы не совсем корректно.

Когнитивную науку, кстати, я называю когнитологией.

На мой взгляд, это лучше, чем когнитивистика, более благозвучно.

Я пригласил к написанию книги «Прологомены когнитивной безопасности» и нейробиолога из Военно-медицинской академии, и специалиста в области когнитивной психологии и когнитивной лингвистики из Санкт-Петербургского госуниверситета.

Такая коллективная монография была издана в 2023 году.

Почему «Прологомены»?

Дело в том, что это название (от греч. *προλεγόμενα* – прежде сказанное) означает некие предварительные замечания, введение в новую область научного знания, ознакомление с ее методами и задачами.

Поэтому нашей целью было издание одного из первых в России комплексных исследований по проблемам когнитивной безопасности в контексте когнитологии как области междисциплинарных исследований, объединяющей творческий поиск философов и психологов, нейробиологов и филологов, социальных антропологов и представителей технических и физико-математических наук.

По крайней мере, мы заявили о том, что проблема когнитивной безопасности достаточно актуально звучит в настоящее время.

Почему?

Дело в том, что при рассмотрении проблем когнитивной безопасности я опирался на отчеты Давосского клуба «Глобальные риски».

Начиная с 2006 года по 2025 год вышло два десятка этих изданий.

Я обратил внимание на то, что еще в предыдущем, 2024 году в отчете было указано, что, по сути, из группы глобальных рисков первое место стала занимать проблема ложной информации и дезинформации.

То есть эта область захватила, условно говоря, всю «поляну» глобальных рисков.

В связи с этим в одном из изданий я рассмотрел проблемы асфатроники, то есть науки о глобальной безопасности.

С чего все началось?

Почему мы сейчас в полный голос начинаем говорить о том, что когнитивная безопасность – это самостоятельное направление глобальной безопасности в целом, национальной безопасности того или иного государства, в данном случае нашего, Российского государства?

Подтверждением этому является ряд документов, в которых эта проблема так или иначе упоминается уже не косвенно, а прямым текстом.

Я даже хотел бы в двух словах напомнить о том, что в прошлом году руководители государств – членов ОДКБ подписали декларацию об объединении усилий в области разработки искусственного интеллекта, где одним из пунктов было записано, что необходимо объединять

усилия в области обеспечения информационной и когнитивной безопасности.

То есть понятие «когнитивная безопасность» все чаще начинает употребляться уже в документах государственного уровня.

Комплексные когнитивные исследования и разработка когнитивных технологий пришлось на вторую половину XX века.

Джордж Миллер как раз описывал события, связанные с научной межотраслевой, междисциплинарной конференцией в Массачусетском технологическом университете, состоявшейся осенью 1956 года.

Летом того же года на Дартмутской конференции, собравшей ряд будущих отцов-основателей ИИ (таких как К. Шеннон, Г. Саймон, М. Мински и другие), по инициативе Джона Маккарти был введен в научный оборот термин «искусственный интеллект».

Формирование когнитивного прорыва сложилось достаточно четко именно с середины 1950-х годов.

Недаром академик РАН Владислав Александрович Лекторский вспоминал о том, что признание в Советском Союзе кибернетики в конце 50-х годов прошлого века инициировало развитие когнитивной психологии, когнитивной лингвистики и, что примечательно, обеспечило своеобразный «когнитивный поворот» в отечественной философии, который совпал с интенсивным развертыванием исследований познавательных процессов⁴.

То есть мировоззренческая проблематика и проблемы, связанные с внедрением кибернетической методологии, объединения ее с естественными и гуманитарными науками, стали явью.

⁴ Лекторский В.А. О прошлом и настоящем. Воспоминания и размышления. М.; СПб., 2025. С. 315-316, 329

И с той же поры, то есть с середины 1950-х годов, и психология когнитивная, и лингвистика когнитивная, и неврология, не говоря уже о кибернетике, стали направляющими трендами развития представлений о когнитологии в целом и постепенно – о когнитивной безопасности в частности.

В 2022 году вышла моя книга «Асфатроника», в которой я обосновал идею и проблемы когнитивной глобальной безопасности.

Приживется или нет такое мое название науки о глобальной безопасности, это уже покажет история. В этой работе я отметил актуальность комплексного исследования глобальной безопасности в следующих ракурсах:

- новая научно-технологическая картина мира включает знания о безопасности существования, развития и функционирования биосферы, социосферы и техносферы, которым присущи общие закономерности и решения на наноуровне (НБИКС-технологии – нано-, био-, информационные, когнитивные и социогуманитарные технологии);

- необходимость обоснования теоретического статуса представлений о феномене безопасности, который до сих пор разнесен между теорией мировой политики (национальная и международная безопасность), безопасностью жизнедеятельности человека, информационной и когнитивной безопасностью;

- ответ на данный вопрос – признание асфатроники как нового направления глобалистики, предметом которой выступает глобальная безопасность, отвечающая на вызовы, связанные с наступлением антропоцена и промышленной революции 4.0.

Наконец-то в 2021 году в обновленную номенклатуру научных специальностей Минобрнауки России

включили новую группу научных специальностей 5.12 – когнитивные науки, в которые вошли следующие научные специальности:

- междисциплинарные исследования когнитивных процессов;

- междисциплинарные исследования мозга;

- междисциплинарные исследования языка;

- когнитивное моделирование.

Причем уже, насколько я знаю, имеются проекты паспортов этих специальностей, создаются диссертационные советы.

По крайней мере, в Институте философии РАН такой совет уже существует, и еще в нескольких городах – по-моему, в Челябинске и Нижнем Новгороде.

Определено, по каким отраслям науки можно защищать кандидатские и докторские диссертации в рамках этих специальностей.

И для меня, честно говоря, непонятно, как это обошло политические науки.

Может быть, со временем жизнь подскажет, что кроме философских, психологических, всего комплекса биологических, естественных наук, технических и физико-математических наук политические науки тоже найдут в этом перечне свое место.

На мой взгляд, это дело будущего.

По крайней мере, уже есть признание.

Это признание и огромного вклада в развитие современной когнитологии, который был сделан нашими отечественными, российскими и советскими учеными.

Я не буду перечислять все фамилии психологов, филологов, философов.

Но важен уже сам по себе феномен того, что достаточно четко была обозначена важность и актуальность проблем исследования связи сознания и психики человека. Об этом даже Джордж Миллер писал в своей статье 2003 года, где говорил о значимости учения И.П. Павлова как родоначальника первой революции в психологии и А.Р. Лурии, который «стал рассматривать мозг и психику как единое целое»⁵.

Это остается фактом, и действительно, нам надо сейчас достаточно активно решать задачу подготовки нового поколения ученых, молодых ученых в области когнитологии.

Вчера просмотрел в видеозаписи выступление Константина Анохина на одной из конференций конца прошлого года, где он, подводя итоги 2024 года в области нейропсихологии, сказал: «Более 500 тысяч статей по всему миру, в основном в европейских университетах, было опубликовано только лишь в достаточно узкой области исследования мозга человека».

И заключает, что, в принципе, это все необходимо освоить, прочитать.

Почему я так заострил внимание на рассмотрении вопроса о когнитивной безопасности?

Дело в том, что начиная с 2020 года появилась серия публикаций, докладов, выступлений на конференциях – в основном натовских военных специалистов, научных специалистов – на тему «Когнитивная война».

В частности, в 2021 году опубликованы материалы проходившей во Франции Первой научной встречи НАТО по когнитивной войне, где четко определялось место

⁵ George A. Miller. The cognitive revolution: a historical perspective // TRENDS in Cognitive Sciences. 2003. Vol. 7. No. 3. P. 143.

когнитивной операции в сфере военных действий как шестой домен, шестое направление:

- когнитивная война рассматривается теперь как самостоятельный, шестой домен в современной войне;

- наряду с четырьмя военными доменами, определяемыми их окружением (наземный, морской, воздушный, космический) и киберсферой, которая их все соединяет, недавние события, нарушающие геополитический баланс сил, показали, как возникла и использовалась эта новая сфера военных действий;

- она работает на глобальной арене, поскольку человечество в целом сейчас подключено к цифровым технологиям и использует информационные технологии и сопутствующие им инструменты, машины, сети и системы.

Это уже даже не чисто информационная война, не информационно-психологическая, а именно когнитивная, когда целью этих операций являются мозг и подсознание человека.

Я нашел один интересный документ, он открытый, опубликован на сайте – это магистерская диссертация 2008 года в одном из военных университетов США.

Лейтенант – сейчас он, наверное, уже постарше в звании – публикует магистерскую диссертацию объемом 300 страниц под названием «Когнитивная война».

Причем он провел анализ на базе арабо-израильского конфликта 1990-х годов, все по полочкам разложил, как надо вести когнитивные операции.

Когда я предложил посмотреть эту работу своим студентам в политехе, сделать обзор разделов, они схватились за голову.

А я считаю, что именно так надо писать магистерские диссертации.

Уже в 2023 году выходит не менее важное исследование подразделения НАТО под названием «Смягчение последствий когнитивной войны».

Это отчет исследовательской группы научно-технической организации НАТО по человеческим факторам и медицине, дающий оценку «науки и технологий, необходимых для смягчения последствий когнитивной войны и защиты от нее».

Отчет был представлен руководству НАТО:

- Cognitive Warfare охватывает широкий спектр передовых технологий (ИИ, ИКТ, машинное обучение) с достижениями в области нейробиологии, биотехнологии и совершенствования человека (аугментации);
- эти технологии разрабатываются в первую очередь на основе достижений когнитивной нейронауки, социологии и культурологии и реализуются в так называемом оперативном цикле принятия решений «Наблюдай, ориентируйся, решай и действуй» (Observe, Orient, Decide, and Act, OODA) с целью сдерживания «противников НАТО» в боевом пространстве XXI века⁶.

То есть произошла своего рода инверсия. Поначалу это были военные наставления для НАТО о том, как надо вести когнитивные войны, а в 2023 году эти специалисты уже представляют ситуацию иначе.

Оказывается, информационная война ведется по отношению к ним со стороны России, им надо защищаться, им надо смягчать эти последствия.

Обращаю внимание на один из интересных разделов этого отчета, достаточно объемного, подробного, со схемами, рисунками, графиками.

⁶ Mitigating and Responding to Cognitive Warfare. Edited by: Dr. Yvonne R. Masakowski (USA) and Dr. Janet M. Blatny (NOR). STO/NATO, 2023.

Столпы, на которых должна строиться процедура смягчения, – это когнитивные нейронауки, социология и культурология.

Эти три столпа должны обеспечивать противодействие когнитивным операциям со стороны России.

Гуманитарные технологии, таким образом, следует рассматривать как одну из важнейших составляющих современных нанобиокогнотехнологий.

ТЕХНОЛОГИИ ОБЕСПЕЧЕНИЯ КОГНИТИВНОЙ БЕЗОПАСНОСТИ

- Философия – разработка социокультурных технологий формирования общественного сознания, идеологии, соответствующей национальным интересам государства и направленной на противодействие деструктивной идеологии.

- Когнитивная нейробиология – применение конкретных исследований защитных структур и функций нейронов и нейронных сетей, управление подсознанием.

- Когнитивная психология в контексте междисциплинарной области исследования сознания, психических процессов (естественно-природной основы развития и функционирования сознания) и социальной деятельности человека нацелена на «принудительное управление» сознанием и, соответственно, обеспечение когнитивной безопасности.

- Когнитивная лингвистика – технологии нейролингвистического программирования и противодействия языковой манипуляции, целью которой является неявное внесение в психику адресата чуждых ценностей, желаний, целей и так далее.

- Когнитивная антропология – интерпретация знаний, традиций представителей различных культур и когнитивных структур, регламентирующих поведение, образ жизни людей, их восприятие мира.

- Когнитивная информатика – регламентация развития искусственного интеллекта и больших данных.

Все направления когнитивной науки имеют свой ход развития, реализации, воплощения, практическое применение в том, что это реализация, формирование когнитивных и нейробиологических технологий.

А также тех технологий, которые сформировались и в когнитивной психологии, когнитивной лингвистике, когнитивной антропологии и информатике (я так называю раздел кибернетического знания и области, связанной с большими данными).

Подводя итог своему выступлению, хотел бы обратить внимание на следующее.

Мы сейчас предпринимаем третий шаг в нашей совместной деятельности, причем наша группа официально никак не организована, да это и не наша задача.

По крайней мере, у нас сложился авторский коллектив, причем охватывающий специалистов не только из Санкт-Петербурга, но и из Екатеринбурга, Читы, Москвы.

В этом году мы издали первый учебник, ориентированный на средства массовой информации, под названием «Когнитивная безопасность человека в медиапространстве».

Насколько он будет полезен, судить не нам.

Я думаю, что со временем такие программы появятся во многих наших вузах.

О необходимости вести такую работу в области предотвращения информационного негативного воздействия сказано в одном из документов по национальной безопасности 2021 года.

То есть наша задача, как мы ее рассматриваем коллективно, – обеспечить систему образования таким учебным материалом.

Сейчас мы будем делать еще одну работу по вопросам преподавания когнитивной безопасности в сфере международных отношений и государственной службы.

Вот такие задачи мы ставим перед собой, и я думаю, что это пойдет только на благо нашего народа.

Спасибо.

Тосунян Г.А.: Спасибо, Игорь Федорович, за содержательный, интересный доклад.

Давайте перейдем к вопросам.

Первым просит слово Кржевов Владимир Сергеевич из МГУ им. М.В. Ломоносова. Пожалуйста, Владимир Сергеевич.

к. филос. н. **КРЖЕВОВ В.С.*** – д. филос. н. **КЕФЕЛИ И.Ф.**

Кржевов В.С.: Благодарю за обращение к очень важной и актуальной теме.

Насколько можно судить, ключевой категорией информационной, когнитивной безопасности является категория истины.

А проблема истины – это классическая проблема философии.

Если вы не обладаете уверенностью, что ваши знания истинны, вы не можете решить никакой связанной с когнитивной безопасностью проблемы.

Маленький пример.

Было сказано о том, что нужно вырабатывать идеологию, отвечающую интересам государства.

Но, хотя интересы государства вполне объективны, они тем не менее не есть нечто самоочевидное.

Они должны быть распознаны.

История дает множество примеров, когда люди действовали в уверенности, что они реализуют свои интересы, а на самом деле они вредили и себе, и другим.

Как быть со всеми этими проблемами в свете концепции информационной или когнитивной безопасности?

Спасибо.

Кефели И.Ф.: Вопрос серьезный, глубокий насчет того, что в когнитивной безопасности важен критерий истины.

А критерием истины является что?

Практика.

Практика и социальная, и идеологическая, а также практика, связанная с научными исследованиями, получением достоверного знания.

* Кржевов Владимир Сергеевич – кандидат философских наук, доцент.

Поэтому, я так полагаю, здесь заострять и абсолютизировать внимание нужно вот на чем.

Истина всегда и конкретна, и относительна, и абсолютна – смотря какую часть этой истинности брать за основу, за исходный спасительный круг.

Я лучше говорил бы о том, что критерием когнитивной безопасности в данном случае является, скорее всего, категория идеального.

Идеальное, понимаемое как то, что является связующим между непосредственной человеческой деятельностью, коллективной деятельностью, и формированием представлений о том духовном мире, который присущ каждому отдельно взятому человеку.

А по отношению к истине ...

Какие мы можем назвать критерии, перечислить на пальцах или как угодно описать абсолютную истину?

Бога нет, или Бог есть?

Что мы охарактеризуем как относительные истины?

Это знания, которые достигнуты на современном этапе.

Всё.

Дальше когнитивная психология открыла то-то, то-то, то-то – и мы уже переходим на новый этап.

Тосунян Г.А.: Спасибо, Игорь Федорович.

Пожалуйста, Галина Вениаминовна. Представьтесь, пожалуйста.

д. филос. н. СОРИНА Г.В. – д. филос. н. КЕФЕЛИ И.Ф.

СОРИНА Г.В.

д. филос. н., проф., профессор кафедры философии языка и коммуникации философского факультета
МГУ им. М.В. Ломоносова.

Сорина Г.В.: Добрый день!

Галина Вениаминовна Сорина, профессор философского факультета Московского государственного университета, кафедра философии языка и коммуникации.

Первое, что я хотела сказать: большое спасибо за приглашение, я первый раз принимаю участие в работе вашего семинара.

Во-вторых, хотела бы поблагодарить Игоря Федоровича за доклад.

В-третьих, прежде чем задать вопрос, отмечу, что с Владимиром Сергеевичем мы не договаривались, но у нас пересекающиеся вопросы.

И я попробую сформулировать вопрос, который уже был затронут, но в какой-то степени данный вопрос требует уточнения и расширения своей явной предпосылки.

В качестве рассуждения – может ли быть информация ложной?

На мой взгляд, информация может быть ложной только на обыденном уровне.

На уровне теоретического знания и теоретического анализа информация ложной быть не может.

Что же может быть ложным или истинным?

Свойством быть ложным или истинным обладает все-таки знание.

И именно на уровне знания мы можем говорить о том, что некоторый поток данных (можно долго спорить о том, как определять информацию) представляет собой,

условно говоря, ложь или истину, адекватность или неадекватность.

Здесь, в общем-то, не идет речь о классике.

В таком случае возникает вопрос: как, с Вашей точки зрения, соотносятся между собой эти два понятия – «информация» и «знание»?

Спасибо.

Кефели И.Ф.: Спасибо.

Информация и знания? Да, это вопрос, который требует сопоставления нескольких, на мой взгляд, подходов.

Мы задаемся вопросом: ложное или истинное знание?

Ложное – это искусственно организованное знание.

Заранее, может быть, перед автором этой ложной истины поставлены цели, задачи.

Я в начале своего выступления уже упомянул данные Давосского клуба.

В отчете по глобальным рискам прошлого, изданном в 2024 году, впервые было заявлено, что ложное знание, ложная информация, которые они не дифференцируют одно от другого, являются наиболее опасным риском.

В группе из десяти рисков, среди которых и продовольственные, и энергетические, и геополитические, по результатам исследований, сбора информации рецензентов, они выделили ложную информацию как наиболее опасный риск в современном мире.

Можно ли этому доверять? Можно ли считать, что это ложь о лжи, либо это все-таки достоверное знание?

Можно спорить, насколько это достоверно, насколько это так или нет.

Поэтому повторяю.

Ложное знание либо является продуктом неполноты исследований, но все равно с преднамеренной

формулировкой выводов, не соответствующих действительности, либо это умышленное препарирование достоверного знания, достоверной информации.

Потому что информация, в ее цифровом варианте, она покрывает все что угодно.

А знание – это уже то, что может артикулироваться человеком на разных языках мира.

Спасибо.

Тосунян Г.А.: Спасибо.

Галина Вениаминовна, можно, я Вам вопрос задам?

Сорина Г.В.: Конечно, можно.

Тосунян Г.А.: Он, возможно, будет дилетантским.

Раз Вы так категорично утверждаете, что информация не может быть ложной, значит, Вы имеете на то достаточно оснований, разделяя информацию и знание.

Если я сообщаю своим друзьям, что сейчас в Москве идет снег и на улице минус 20 градусов, я сообщаю им информацию, а не передаю им знания.

Эта информация является ложной. Сказать, что я передаю им ложные знания, было бы, наверное, некорректным. Объясните мне это противоречие.

Сорина Г.В.: Большое спасибо за вопрос.

Я бы ответила следующим образом.

Если Вы сейчас своим друзьям говорите то, что Вы сказали, то Вы их информируете о том знании, которым Вы обладаете.

А если это знание Вами сознательно сформулировано ложным образом, то Вы преследуете определенную цель.

И тогда возникают другие вопросы: какова эта цель, зачем и так далее.

Проблема знания, в отличие от информации, на мой взгляд, связана с тем, что знание – это все та же совокупность данных, прошедших через нас и присвоенных нами, когда мы можем ими оперировать.

Вы можете нечто знать.

Но когда Вы распространяете имеющееся у Вас знание, то возникает еще другая, дополнительная проблема – это вопрос об отношении к Вашему знанию.

Он предстает как проблема доверия к Вашему знанию и к Вам как распространителю этого знания.

А поток информации, нерасчлененный поток данных, который идет, нельзя формулировать как истинный или ложный.

Для того чтобы выявить это отличие, Вам, на мой взгляд, нужно иметь критерии, а критерии – это знание.

В общем, это отдельная проблема.

Тосунян Г.А.: Да, это отдельная проблема. Но все-таки Ваше категоричное утверждение немного режет слух.

Если Вы серьезно занимались этим вопросом, подумайте, может быть, сделаете специальное сообщение на эту тему на одном из наших заседаний.

Сорина Г.В.: Спасибо большое.

Тосунян Г.А.: Спасибо, Галина Вениаминовна.
Виктор Дмитриевич Кривов, пожалуйста.

д. э. н. КРИВОВ В.Д. – д. филос. н. КЕФЕЛИ И.Ф.

КРИВОВ В.Д.

д. э. н., заведующий кафедрой статистики экономического факультета и заведующий кафедрой социологии и менеджмента общественных процессов Высшей школы современных социальных наук МГУ им. М.В. Ломоносова

Кривов В.Д.: Спасибо, уважаемый Гарегин Ашотович, уважаемые коллеги.

Спасибо, Игорь Федорович, за этот очень интересный доклад.

У меня вопрос.

У Вас упоминался термин «асфоцефатроника».

А в каком смысле Вы его применяете к концепции когнитивной безопасности?

Спасибо.

Кефели И.Ф.: Спасибо. «Асфолей» – это безопасность, до сих пор этот термин сохранился в греческом языке.

Асфатроникой я называю определение глобальной безопасности.

А асфоцефатроника – это уже наука о безопасности собственно когнитивной.

Цефализация, цефал – это в переводе с греческого мозг, голова.

Поэтому термин «асфоцефатроника» я использовал, исходя из этого.

Кривов В.Д.: То есть это синоним когнитивной безопасности в Ваших исследованиях?

Кефели И.Ф.: Когнитивная безопасность – это область действий, технологий.

А асфоцефатроника – это наука о когнитивной безопасности, объединяющая и собственно фундаментальные исследования, и прикладные направления применения когнитивных технологий, которые так или иначе связаны с интеллектуальной деятельностью, с функцией мозга и языка и так далее.

Кривов В.Д.: Спасибо.

Тосунян Г.А.: Спасибо.

Сергей Федорович Сергеев. Представьтесь, пожалуйста.

д. п. н. СЕРГЕЕВ С.Ф. – д. филос. н. КЕФЕЛИ И.Ф.

СЕРГЕЕВ С.Ф.

д. п. н., профессор кафедры информационных систем
в искусстве и гуманитарных науках
Санкт-Петербургского государственного университета

Сергеев С.Ф.: Добрый день, уважаемые коллеги, уважаемый председатель!

Позвольте представиться: профессор Санкт-Петербургского университета, заведующий лабораторией сложных систем Санкт-Петербургского политехнического университета Петра Великого, председатель Санкт-Петербургского отделения научного совета по методологии искусственного интеллекта и когнитивных исследований при Президиуме РАН.

Хочу прежде всего поблагодарить глубокоуважаемого Гарегина Ашотовича за любезное приглашение на данный семинар и предоставленную мне возможность выступить.

Семинар чрезвычайно интересный и содержательный, как и выступление Игоря Федоровича Кефели.

Оно сделано прекрасно, и хотелось бы в порядке дискуссии дать критику лишь некоторым его положениям.

Действительно, когда речь идет об информационной безопасности, то всегда возникает вопрос: безопасности от какой информации и по какой причине?

В этом плане мне понравился комментарий Галины Вениаминовны Сориной, в котором она четко отметила, что знание и информация – это совершенно разные вещи.

Я поддерживаю данный тезис и хотел бы раскрыть его подробнее.

Знание есть присвоенная субъектом и включенная в его когнитивные инструменты и структуры, связанная в

единое семантическое целое информация, позволяющая субъекту непротиворечиво познавать мир.

Это индивидуальное структурно-семантическое приобретение человека, неотделимое от его когнитивной структуры и работающее с сознанием.

Знание – это форма существования системы «организм – окружающая среда».

Оно проявляется в процессе реорганизации системы в действиях ученика при достижении им требуемого результата.

При этом нет никакой передачи знаний.

Как следствие, знание нельзя отделить от его носителя и среды. Оно, будучи распределено в системе «организм – окружающая среда», является свойством, а не продуктом системы.

Являясь аутопоэтической, самоорганизующейся, операционально замкнутой системой, наш мозг не пропускает в зону конструирования субъективного мира никакую информацию, которая может в той или иной мере разрушить аутопоэзис сознания – сохранение простой, ясной и устойчивой картины мира, его понимания и расширения в деятельности и так далее.

Мозг строит и защищает эту картину всеми возможными способами.

Часто говорят, и здесь прозвучало, что информация может что-то сделать с человеком.

Нет – непосредственно с мозгом ничего не может сделать! Он не пропускает в нашу психику никакую информацию, которая может каким-то образом разрушить ясность и целостность нашего сознания.

Однако нас можно легко ориентировать в нужных кому-то (в том числе и негативных) направлениях и манипулировать нами.

Это первое.

Второе.

Когда речь идет о когнитивных науках, то здесь мы начинаем иметь дело с легкой спекуляцией.

Спекуляцией в каком плане?

Дело в том, что Джордж Миллер и его коллеги на известном симпозиуме, состоявшемся в Массачусетском технологическом институте (MIT) в 1956 году, попробовали конституировать когнитивную науку как область междисциплинарных исследований познания.

Но это был лозунг для объединения естественных и гуманитарных наук, а вовсе не новая наука.

И, собственно, жизнь показала, что конвергенции дисциплинарных областей психологии, физиологии и компьютерных наук в целое не произошло, несмотря на достигнутый прогресс в каждой из них.

Когнитивная наука как отдельная наука так и не возникла.

Это по-прежнему совокупность наук о познании.

Вместе с тем мы видим попытки ее создать в нашей стране хотя бы формально.

Так, Высшая школа экономики в 2021 году ввела в номенклатуру научных специальностей ВАК группу специальностей 5.12 – когнитивные науки.

Мы сейчас видим, что за это время в их совете защитились всего лишь один доктор когнитивных наук – Василий Андреевич Ключарёв (в 2024 году) и три кандидата наук по этой же дисциплине.

То есть, повторяю, попытка создания в процессе междисциплинарного синтеза конвергентного ядра психологии, философии, лингвистики, искусственного интеллекта, антропологии и неврологии работает плохо.

Если вы внимательно посмотрите на эти работы, то поймете их технологию и принцип построения исследований.

Берется какой-либо высокотехнологичный метод исследования деятельности мозга, например, магнитно-резонансная томография (МРТ), и метод воздействия на него, например, транскраниальная магнитная стимуляция (ТМС). И с помощью анализа данных, полученных в некоторой тестовой ситуации, делаются выводы о психологических свойствах человека – например, уровне его конформизма.

Или с помощью ТМС пытаются изменить поведение человека при принятии решений.

Это суть исследований в рамках «нейроэкономики».

Получают много трудноинтерпретируемых данных, но пока никаких серьезных теоретических и практических достижений в этой области, кроме красивых названий, нет.

Та же «опера» и с другими достижениями когнитивных наук.

Перейдем к нашему обсуждению и посмотрим, что такое когнитивная безопасность.

Безопасность когнитивных систем человека, систем, которые познают мир...

Ведь когнитивная система, по определению, это система, познающая мир.

«Когнито» – от латинского *cognitio* («познание»).

Познавать мир могут все биологические живые системы в процессе своей эволюции.

Это суть жизни. Познание – это жизнь!

Мне понравилось, как точно сказал Игорь Федорович о когнитивной безопасности: это защита от манипуляций, защита от каких-то действий враждебных сил.

То есть мы начинаем понимать, что комплекс проблем, о котором мы говорим, не является чистой наукой.

Скорее, это комплекс мер, включающий разные способы ограничений информационного и иного злонамеренного воздействия на ситуации и ориентацию субъекта.

Это такая попытка объединить на практике несоединимое в рамках научного метода.

А вопрос у меня к докладчику простой.

В чем заключается «когнитивный поворот», и что такое, по его мнению, когнитивные науки?

Что обозначают прозвучавшие в докладе термины «когнито», «когнитология»?

Тосунян Г.А.: Спасибо. Пожалуйста, Игорь Федорович.

Кефели И.Ф.: Из сказанного Сергеем Федоровичем я понял некоторые опасения с его стороны, что когнитология – это не наука, или когнитивные науки – это не науки, а просто набор отдельно взятых дисциплин.

Но я об этом не говорил.

Я говорил о том, что это, по сути дела, междисциплинарная область.

Джордж Миллер написал об этом в небольшой статье, уже будучи в преклонном возрасте, в 2003 году.

Он объяснил, что тогда думал, что когнитивная междисциплинарная область сформировалась, сложилась.

В то время кибернетики оставались кибернетиками, психологи оставались психологами и работали в своих сферах.

Но сам тренд, связанный с когнитивным поворотом, в первую очередь с зарождением когнитивной психологии, – это факт исторический.

И когнитивный поворот уже потом разошелся по всем направлениям.

То есть центр внимания – эта связка: сознание и мышление, сознание и язык.

Миллер даже шестигранник нарисовал в этой статье 2003 года и обозначил, что линии, связующие вершины этого шестиугольника, как раз и являются областями междисциплинарных исследований.

А то, что Вы говорите о когнитивном повороте и когнитивной психологии...

Психология, повторю, это одно из направлений, зародивших когнитологию, тесно с ней связанных.

В центре когнитивного поворота были философские рассуждения, связанные с проблемой истины, с проблемой лжи, достоверного знания, которые мы слышим и воспроизводим до сих пор.

Поэтому исследования, связанные с когнитивной безопасностью, заключаются еще в том, что негативное влияние может оказываться не только на сознание человека, и это влияние он может на себе даже и не наблюдать, но и более глубоко, на подсознание.

Некоторые выводы об этом уже можно услышать, они даже публикуются. Может быть, в силу своего секретного характера они пока широко не обсуждаются, но это факт.

Подсознание – это то, что присуще человеку.

Сознание – у бодрствующего человека, а подсознание работает от рождения до его последнего вздоха на земле.

И сейчас подсознательная область тоже становится предметом воздействия определенных когнитивных атак.

Спасибо.

Тосунян Г.А.: Спасибо. Коллеги, если не возражаете, заслушаем второй доклад, потом зададим вопросы второму докладчику и перейдем к обсуждению.

У нас высказали желание принять участие в обсуждении Конявский Валерий Аркадьевич – заведующий кафедрой защиты информации МФТИ; Гузов Юрий Николаевич – доцент, научный руководитель магистерской программы из Санкт-Петербурга; есть и еще желающие выступить.

Но сначала мы заслушаем второго докладчика.

Михаил Олегович Колбанёв, пожалуйста, Вам слово.

ДОКЛАД 2

КОЛБАНЁВ М.О.

д. т. н., проф., профессор кафедры информационных систем
и технологий Санкт-Петербургского государственного
экономического университета

ИСКУССТВЕННЫЙ ИНТЕЛЛЕКТ КАК ХАЙП

Здравствуйте, уважаемые коллеги, уважаемый Гарегин Ашотович! Большое спасибо за предоставленную возможность выступить перед таким высоким собранием.

Из названия моего доклада следует, что объектом исследования является искусственный интеллект.

Хочу подтвердить слова нашего ведущего, что искусственный интеллект – это, действительно, многозначный термин.

И разные люди понимают его по-разному.

Но еще хуже, что и разные специалисты понимают этот термин по-разному.

Я намеренно дал своему докладу провокационное название.

Я использую слово «хайп», для того чтобы обострить соответствующую проблему.

В качестве цели сообщения я вижу развитие системного взгляда на полезности и опасности особых компьютерных технологий, которые называют «искусственный интеллект».

Попробую описать две линии, которые можно обнаружить в развитии искусственного интеллекта (ИИ):

– с одной стороны, ИИ – это притча во языцех, объект пристального общественного внимания, окруженный хайпом, а также цель как для насмешек над его

«невиданными возможностями», так и для реальной деятельности, декларируемой некоторыми акторами;

– с другой стороны, ИИ – это цифровая технология, характеризующаяся определенными физическими, энергетическими, информационными, экономическими и другими показателями.

Само сообщение содержит две части:

1. В первой из них предлагается эмоциональный взгляд на искусственный интеллект, отражающий особенности хайпа в технологической сфере вообще и выделяющий хайп «за» и «против» искусственного интеллекта в частности.

2. Во второй части развивается рациональный взгляд на искусственный интеллект (антихайп). Здесь речь идет об ИИ в контексте цифровой трансформации, архитектуры ИИ, его экономического портрета и, наконец, об ограничениях ИИ.

1. Эмоциональный взгляд на искусственный интеллект (хайп)

Искусственный интеллект – многозначный термин, имеющий целое поле смысловых значений:

– от самого широкого (ИИ – это высший уровень развития цифровых технологий),

– до самого узкого (ИИ – это нейросетевая технология генерации контента).

Такой разброс приводит к противоречивым мнениям о том:

– в каких областях деятельности использование систем ИИ будет эффективным;

– кто должен стать основным потребителем услуг ИИ;

- какие технологии ИИ надо развивать в приоритетном порядке;
- как обеспечить защиту от опасностей, которые создаются ИИ, и о многом другом.

Именно отсутствие общепринятого понятийного аппарата дает простор для хайпа в области ИИ.

Я понимаю под хайпом действия, направленные на обман, дезинформирование кого-либо путем информационной шумихи (по принципу «если соврать много раз, то тебе в конце концов поверят»).

Хайп, как я это понимаю, всегда имеет две стороны:

- недобросовестный, лживый источник хайпа;
- доверчивая, наивная жертва хайпа.

Таким простым примером хайпа, совершенно безобидным, может быть типичная для многих бытовая ситуация.

Группа людей на берегу моря может начать шумиху о том, что там плывут кит или дельфины, хотя их там нет.

И многие в это поверят.

Не поверят только те, кто знает о том, что в этой акватории не может быть китов или дельфинов.

Но людей, которые понимают это, обычно в каждой конкретной ситуации всегда намного меньше, чем тех, которые не обладают таким знанием.

Аналогично, если завалить прессу сообщениями о небывалых возможностях искусственного интеллекта, то в это многие могут поверить, хотя большинство из этих сообщений могут быть недостоверными.

Опасность хайпа как формы обмана нельзя уменьшать.

В 2025 году Давосский ВЭФ отнес дезинформацию к одной из трех ключевых угроз человечеству, так как она разрушает информационное пространство и нарушает базовые права человека.

Важные элементы схемы обмана про ИИ – это замалчивание правды о его ограничениях и распространение неправды о его возможностях.

В общем случае хайп в области ИИ можно разделить на три вида:

- хайп «за» ИИ,
- хайп «против» ИИ,
- антихайп.

Хайп «за» искусственный интеллект использует уже апробированную не один раз схему дезинформации в сфере технологий на глобальном уровне.

В качестве примера можно выделить три хайповые технологии.

Тема искусственного интеллекта пришла к нам вслед за шумихой вокруг борьбы с технологиями, якобы изменяющими климат Земли, а она в свою очередь сменила «борьбу» за восстановление озоновой защиты Земли.

Вехи распространения всемирных информационных цунами, продвигающих эти концепции, имеют ряд сходств:

- сначала ученые строят статистические модели очень сложных природных процессов: озонового слоя, климата Земли и мышления человека, соответственно;

- затем допускается истинность этих моделей (хайли лайкли); более того, авторам вручаются Нобелевские премии в области физики;

- вишенкой на этом торте из статистических пузырей становятся масштабные конференции в Вене (охрана озонового слоя) и Париже, на которых назначается материальная ответственность за использование сначала фреона, затем углеводородного сырья и, наконец, премия за обеспечение в недалеком будущем беспрецедентного повышения эффективности деятельности с помощью ИИ.

Разбогатеть от внедрения альтернативных технологических проектов всегда стремились страны и предприятия Запада за счет стран, использующих свою, отличную от западной, технологию.

А кто назначен бенефициаром, и кто должен остаться в дураках из-за дезинформации в области цифровых технологий нового типа?

Примерами недобросовестной шумихи вокруг ИИ могут служить следующие заявления, которые аргументированно опровергаются во множестве публикаций:

- искусственные нейронные сети подобны мозгу человека (это утверждение встречается в огромном числе публикаций). Это не так. Например, мозг человека требует ~ 1 Вт·час, а только для обучения «хорошей» нейросети надо $10^{25} \div 10^{26}$ флоп или ~ 300 МВт·час энергии;

- гендиректор OpenAI: агенты сильного ИИ (AGI), способного думать, как человек, появятся в 2025 году. Это не так. Интеллект человека и цифровые системы вообще не сопоставимы, ИИ никогда не сможет «думать»;

- технология ИИ способна заменить человека в ряде областей. Это не так. ИИ нигде не сможет работать без человеческого управления и контроля;

- ИИ приносит экономический эффект, повышает производительность бизнеса на 40%, около 55% компаний в мире уже используют ИИ. Это не так. Экономический эффект от внедрения ИИ не очевиден и не доказан. Напротив, в США об использовании ИИ заявляют только около 5% компаний. В то же время ИИ является источником многих системных проблем;

- к 2030 году международный рынок ИИ увеличится в двадцать раз. Это не так. Пока что актуален вопрос о выживаемости ИИ-компаний;

– ИИ способен стать Творцом, выполнять творческие функции, а не только компилировать отрывки доступного контента. Это не так. ИИ не может выйти за рамки известного, его максимум – комбинация доступных ему фрагментов данных;

– развитие нейронных сетей неизбежно приведет к замене традиционных методов познания преобразованием цифровых данных. Это не так. Логический способ познания не единственный и не самый главный. В основе нейронных сетей лежит теорема Колмогорова – Арнольда. Реальная жизнь – это не набор датасетов (организованных наборов данных, используемых для обучения, тестирования и валидации моделей машинного обучения, текстов, фото, звуков, лиц и так далее).

Искусственный интеллект окружен сегодня плотной завесой:

– некорректных, неоднозначных, «мутных» сообщений о его возможностях и внедрениях,

– преувеличений об эффекте от его использования в настоящем и будущем,

– наконец, ИИ окутан прямым враньем о его системный свойствах.

Вот еще примеры хайпа из актуальной ленты новостей:

– в Сколково ИИ уже заменил или вот-вот заменит профессора и будет учить еще лучше, чем учит этот профессор. И в этом есть какая-то часть правды. В чем-то искусственный интеллект может помочь профессору. Но если в процессе учебы обучаемый думает, мыслит – тогда, конечно, как тут заменить профессора?;

– или другой, тоже удивительный хайп из нашей ленты новостей. Трамп принял решение о бомбардировке Ирана по совету ИИ. Здесь, может быть, тоже есть часть правды. Конечно, при принятии этого решения

обрабатывались большие объемы данных специальными цифровыми системами;

– или вот еще: Алиса – искусственный интеллект Яндекса – придет в каждый дом и изменит жизненный уклад многих семей.

Увы, в истинность этой недостоверной информации верят многие.

Хочу предложить ключ к хайпу «против» искусственного интеллекта. Хайпа «за» искусственный интеллект очень много. А хайп же можно строить и «против». Хотя, когда «против» чего-то, это тоже негативный процесс.

Я называю эту схему «Оба без ума». Напомню диалог Чацкого и Репетилова из комедии А.С. Грибоедова «Горе от ума»:

Репетилов

... у нас... решительные люди,
Горячих дюжина голов!
Кричим – подумаешь, что сотни голосов!..

Чацкий

Да из чего беснуетесь вы столько?

Репетилов

Шумим, братец, шумим...

Чацкий

Шумите вы? и только?

Репетилов

Не место объяснять теперь и недосуг,
Но государственное дело:
Оно, вот видишь, не созрело,
Нельзя же вдруг.

Комедия была написана 200 лет назад, но тем не менее представьте, что это диалог об искусственном интеллекте!

Могу предположить, что теперь, в XXI веке, Александр Сергеевич мог бы с удовольствием посетить спектакль «Горе без ума» (трагикомедия М. Розовского в театре «У Никитских ворот») и написать трагедию «Оба без ума», в которой:

– «Чацкий» – это бездушный и бездумный нейросетевой ИИ,

– «Репетилов» – это безбашенный и недалекий его поклонник.

Можно себе представить, какие решения будут приняты в результате взаимодействия этих двоих без ума.

А кто же должен думать и принимать решения, если оба без ума?

В основу хайпа «против» искусственного интеллекта можно положить простое и истинное утверждение: ИИ – это просто часть киберпространства, которая не способна думать и заменить человека там, где необходима ответственная, вдумчивая созидательная деятельность.

И 200 лет назад, и сейчас информационный шум, дезинформация находятся на службе у нечестных дельцов и являются способом обмана жалких неудачников, неучей, дураков (таких как Репетилов).

Отделять хайп про ИИ от его реальных возможностей надо при помощи антихайпа, то есть изучения принципов построения, архитектуры и системных свойств технологий ИИ.

2. Рациональный взгляд на искусственный интеллект (антихайп)

С моей точки зрения, говорить об ИИ вне процесса цифровой трансформации – совершенно пустое дело.

А процесс цифровой трансформации неоднороден в своей структуре.

Я выделяю в этом процессе три важных этапа.

Первый называю – цифровизация, второй – цифровая экономика, а третий – экономика данных.

При этом я использую те термины, которые введены государственными документами, программами в нашей стране.

Принципиальные отличия этих этапов заключаются в контексте работы с данными:

1) на этапе цифровизации, или «Кибернетики+», были развиты ручные технологии оцифровки аналоговых процессов;

2) на этапе цифровой экономики, или «Интернет+», было создано киберпространство, в котором накоплен огромный объем цифровых данных, формируемых не людьми, а технологиями.

Достаточно сказать, что в 2025 году создается по 60 гигабайт данных в день на каждого жителя Земли. Это трудно даже вообразить. Это каждому из нас, я так себе это представляю, каждый день кто-то вручает флешку 64 гигабайта, а на следующий день еще одну, а потом еще одну, и говорит: «Используй эти данные для принятия решений, для управления теми или иными процессами».

Понятно, что человеческие возможности не позволяют это сделать;

3) на этапе экономики данных, или «ИИ+», необходимо повысить КПД работы в информационной среде благодаря организации данных на цифровых платформах и другим технологическим решениям.

Таким образом, в процессе цифровой трансформации изменился порядок добывания информации:

– на этапе цифровизации люди сначала добывали информацию, а затем превращали ее в цифровые данные, расширяя технологические возможности сохранения, распространения и переработки. Так, например, бумажный учет получил цифровую форму;

– в цифровой экономике все наоборот. Сначала сквозные технологии без участия людей создают цифровые данные, а затем люди при помощи алгоритмов переработки добывают из них информацию. Такая перемена и создает потребности в создании систем ИИ.

Сложившееся положение дел дает возможность сформулировать следующее определение: искусственный интеллект – это технология добывания информации из данных (массива цифр), основанная на небывалой производительности современных цифровых систем.

Для практических задач такой обобщенный взгляд должен быть дифференцирован. Выделю три разных, но не противоречащих друг другу подхода к пониманию ИИ, которые отражают ответственность государства, бизнеса и науки.

Для государства ИИ – это новая мощная цифровая технология – экономика данных.

Президент В.В. Путин: «...наши разработки в сфере ИИ занимают лидерские позиции. В том числе в сфере цифровизации мы ведем, входя в десятку лучших на мировом уровне».

Похожим образом высказываются Вэнс, Макрон и другие.

Определения ИИ в государственных документах (национальных проектах, министерских приказах, ГОСТах и так далее) ограничены областью их применения. Однако нельзя не заметить, что определения ИИ формулируются и на высоком государственном уровне.

Известны указы президентов РФ и США (еще Байдена). В ЕС принят закон об ИИ. Там дается такое определение:

«ИИ – компьютерная система, которая:

– предназначена для работы на различных уровнях автономии;

– может проявлять адаптивность после развертывания;

– делает выводы на основе получаемых входных данных о том, как генерировать выходные данные в форме прогноза, **контента**, рекомендации или решения для достижения явных или неявных целей;

– может быть полезной и влиять на физическую или виртуальную среды».

Ровно такое определение искусственного интеллекта, только без понятия контента, мы с моим учителем и коллегой, профессором ЛЭТИ С.Я. Яковлевым дали в книге «Модели и методы оценки характеристик обработки информации в интеллектуальных сетях связи», которая еще в далеком 2002 году была издана в нашем Санкт-Петербургском государственном университете.

Бизнес видит ИИ с других позиций.

Хочу сказать, что для бизнеса терминология имеет второстепенное значение, не такое важное, как в других областях.

Для потребителей, а не разработчиков цифровых технологий важно, чтобы та технология, которую они используют, приносила прибыль, уменьшала, сокращала издержки или имела еще какой-то положительный результат.

И не важно, как она называется.

При таком взгляде ИИ надо рассматривать избирательно, как набор разных технологий.

Например, Ростелеком относит к ИИ следующие технологии:

– машинное обучение,

– компьютерное зрение,

– поиск неструктурированной информации,

– обработка естественных языков,

- технологии распознавания и синтеза текста, лиц, речи, жестов, звуков, запахов и прочего,
- технологии поддержки принятия решений,
- технологии биометрии,
- геоинформационные технологии и навигацию,
- интеллект роя, умную пыль и другие.

Легко заметить, что большинство из этих технологий не имеет отношения к выработке контента как результата работы цифровых систем.

Замечу еще, что Гарегин Ашотович совершенно прав, когда говорит о том, что разработчики технологий называют искусственным интеллектом то, что продают, потому что это дает им какие-то преференции на рынке.

И инвесторы охотнее будут вкладывать деньги, если их вложения называются «искусственный интеллект».

Такая стратегия облегчает бизнесу выход на рынок.

Безусловно, для науки понятийный аппарат имеет первостепенное значение. Иначе будет возникать подмена терминов, и ученые не будут понимать, один и тот же или разные объекты они обсуждают.

Наука призвана формировать знание, которое должно быть понятно всем исследователям.

В научных определениях выделяется три признака технологий ИИ:

1) иное по сравнению с другими технологиями взаимодействие с внешним окружением, способность заменить человека в интеллектуальной деятельности (ИИ – идеальный раб: он делает то, что велят, и при этом абсолютно послушен [Н. Винер]);

2) применение алгоритмов добывания информации, использующих передовые методы работы с цифровыми данными;

3) привлечение входных цифровых данных, которые добыты без участия людей, свободны от влияния человеческого фактора. Согласитесь, что если входные данные технологий дает человек, то он способен через эти входные данные так или иначе влиять на решения. И в этой ситуации трудно говорить о том, что искусственный интеллект заменяет в чем-то человека.

В психологической науке ИИ – это модель интеллекта человека, неприменимая к технической системе.

Проведенный анализ позволяет разделить понимание ИИ на две группы.

В широком смысле ИИ – это любая цифровая технология, способная заменить (или, как говорится в указе президента РФ, «имитировать когнитивные функции») человека в каком-либо процессе деятельности.

Главная опасность ИИ в таком контексте заключается в том, что принятие решений и опыт специалистов формируются на основе изучения цифровых данных, а не натурального эксперимента.

Человек в результате может потерять связь с реальными, природными физическими и социальными процессами.

Всегда существует вопрос: а дают ли доступные цифровые технологии адекватное описание реальности?

В узком смысле ИИ – это цифровые технологии, использующие архитектуру нейронных сетей и методы их обучения; в важном частном случае – это технологии, генерирующие контент.

Главная опасность технологий такого рода заключается в том, что выводы о состоянии процессов и систем являются результатом компиляции данных, предварительно отобранных и размеченных людьми.

Эти данные всегда отражают малую часть опыта и знаний, которыми располагает человечество в соответствующей области знания, они ангажированы.

При этом большая часть людей может принимать полученный контент как истину, поскольку его форма похожа на ту, которую используют специалисты, хотя она и носит не смысловой, а статистический характер.

Промежуточное положение между этими подходами занимают сквозные технологии, которые и формируют, и сохраняют, и преобразуют большие объемы данных.

Во всех случаях ИИ – это часть киберпространства, а не реального мира.

Строго говоря, все цифровые системы являются интеллектуальными.

Самые большие инвестиции вложены в инфраструктуру – цифровые сети и системы хранения.

Первые переносят данные в пространстве и заменяют телефонных барышень, вторые – переносят данные во времени и заменяют библиотекарей и архивариусов.

Компьютерные системы (системы переработки данных) автоматизируют вычисления.

Если понимать ИИ в узком смысле как генератор нового контента на основе того, на котором он был обучен, то такая цифровая система объединяет все те же цифровые сети, системы хранения и вычислители, но используются они для выполнения эмпирических алгоритмов вычисления.

ИИ, генерирующий новый контент, реализует в ходе вычислений особую математическую схему, в основе которой лежат:

- искусственные нейронные сети;

– процедуры статистического обучения этих сетей на данных, выбранных специалистами из информационного пространства и затем определенным образом размеченных метаданными;

– вероятностная процедура формирования по промпту выходных данных из фрагментов контента, сохраненного при обучении.

Чемпионом по инвестициям в такой ИИ является американская компания OpenAI, которая на март 2024 года вложила в разработку генеративного ИИ \$14 млрд,кратно больше, чем любая другая компания в мире.

Это большие деньги, но если сравнить их с общими мировыми расходами на информационные технологии в 2025 году, которые, по прогнозу Gartner, достигнут \$5,6 трлн, то они уже не кажутся такими большими. Это всего 2,5%, то есть 2,5 промилле (тысячных долей), от тех денег, которые мировое сообщество тратит на цифровые технологии.

Представляется важным архитектурный подход к описанию ИИ.

Оттолкнемся от понятия «архитектура сложности», которое ввел нобелевский лауреат Герберт Саймон.

Самая известная и применяемая иерархическая модель такого типа – это описание структуры философского знания, где выделяют: онтологию, гносеологию и аксиологию.

Аналогично в устройстве компьютера выделяются аппаратное обеспечение, программное обеспечение и предметная область, в психофизиологии человека – физиология, психика и поведение, и тому подобное.

Если применить этот подход к системе ИИ, то очевидным образом выделяются:

– на нижнем уровне – сквозные технологии, которые формируют цифровые данные о состоянии физического и социального пространств без участия человека. Повторю, если данные цифровой системе «подсовывают» люди (например, датасеты), то это уже не искусственный, а естественный интеллект, он ангажирован, подлежит контролю;

– на среднем уровне архитектуры находятся алгоритмы систем ИИ, которые можно разделить на две группы. К первой относятся те, для которых доказана применимость математического аппарата при описании соответствующей предметной области, а результаты расчетов объяснимы и интерпретируемы (программирование). Ко второй – применимость алгоритмов для решения задач объясняется правдоподобными рассуждениями о свойствах статистических моделей. Это, в частности, системы нечеткой логики, эволюционные, звериные (имитирующие поведение животных) алгоритмы и тому подобное, а также машинное обучение, нейронные сети, глубокое обучение (промпт-инжиниринг);

– на верхнем прикладном уровне – пока что предположительно – ИИ будет применим к задачам, которые можно решить методами классификации, регрессии, кластеризации и уменьшения размерности, например, в банковском и страховом деле, телекоммуникациях, здравоохранении, производстве, торговле, маркетинге и других сферах. Особое место занимают большие языковые модели, так как именно язык человеческого общения является носителем смыслов.

Приведу статистику с авторитетного сайта по затратам, корпоративным инвестициям в искусственный интеллект.

В 2021 году в ИИ было инвестировано 276 млрд.

Гарегин Ашотович назвал цифру 300 млрд.

Источники разные, но цифры «бьются».

Это, конечно, деньги не на генеративный искусственный интеллект, а на широкий спектр цифровых технологий. Вспомним перечень технологий, который дает Ростелеком.

Однако в последнее время рынок ИИ характеризуется резким уменьшением корпоративных инвестиций и затрат на разработку программного обеспечения.

Уже в 2022 году инвестиции упали на 40% по отношению к 2021 году.

С тех пор инвестиции уменьшаются на 25% год к году. Темпы роста затрат на программное обеспечение за период с 2021 года уменьшились в два раза.

Главные прикладные области, в которых работают указанные инвестиции, можно охарактеризовать следующим образом:

- смартфоны (ИИ применяется в умном видео, голосовом помощнике, для распознавания лиц и так далее);
- социальные сети (например, для персонализации контента);
- интернет вещей (умные вещи, колонки, термостаты и многие другие устройства).

Цели обращения людей к генеративному ИИ сводятся к нескольким главным:

- помощь при ответах на сообщения друзей и коллег;
- решение финансовых задач;
- планирование маршрутов;
- создание электронных писем;
- подготовка к собеседованию.

Системы искусственного интеллекта, построенные на машинном обучении, достаточно неточны в своих прогнозах и оценках.

Машинное обучение предсказывало изменения на фондовом рынке с точностью 62%, смертность пациентов с COVID-19 – 92%, рак молочной железы – 89%, предотвращение кибератак – 86% и так далее.

Это совсем не та точность, к которой привыкли инженеры, создающие традиционные цифровые системы.

Приведу несколько фактов, которые говорят об ограниченности искусственного интеллекта.

Во-первых, искусственный интеллект, особенно генеративный искусственный интеллект, непредсказуем: он непознаваем и для своего создателя, и для пользователя.

ИИ в том виде, как его создают нейронные сети, является черным ящиком.

Этот факт прямо отмечается в указе президента РФ «О развитии искусственного интеллекта (ИИ) в РФ», который в том числе утверждает и Национальную стратегию развития искусственного интеллекта до 2030 года: «Отсутствие понимания того, как ИИ достигает результатов, является одной из причин низкого уровня доверия к современным технологиям ИИ...»

Я хочу обратиться к авторитету академика В.Б. Бетелина, одного из конструктивных критиков искусственного интеллекта.

Он прямо доказывает, что для полуэмпирических зарубежных нейросетевых алгоритмов, доступных в сети, отсутствуют какие-либо теоретические обоснования, они уязвимы для кибератак, и вообще их использование опасно.

Академик указывает: «Природа ошибки заложена в самой сути современных нейросетей».

Вот что говорят другие эксперты:

– Линус Торвалдс, один из самых авторитетных программистов нашего времени: «Сегодня ИИ состоит на

90% из маркетинга и только на 10% из реальности. И вообще, эта шумиха не может не раздражать»;

– глава Baidu (крупнейшей поисковой системы Китая) называет ИИ пузырем и предсказывает банкротство 99% ИИ-компаний.

Недавно журнал «The Economist» опубликовал статью «Лопнет ли пузырь искусственного интеллекта в 2025 году, или он начнет приносить плоды?».

В ней утверждается, что развивать ИИ становится очень дорого:

– расходы на ЦОДы для ИИ в 2024–2027 годах превысят \$1,4 трлн;

– стоимость обучения следующего поколения моделей – \$1 млрд;

– компания OpenAI примерно 75% дохода получает не от компаний, а от физических лиц;

– только 5% компаний в США заявляют об использовании ИИ в своих продуктах и сервисах.

Возникает законный вопрос: зачем OpenAI и другие тратят миллиарды, если бизнес в США не готов платить за эти услуги?

Ответ, на мой взгляд, прост: генеративные технологии создаются для туземцев (или Репетиловых), а для их внедрения за пределами США нужен хайп «за» ИИ.

Где предположительно можно использовать нейросетевой и генеративный искусственный интеллект?

Прежде всего это творчество.

В то время как сам ИИ не способен к творчеству, он может использоваться как новый уникальный инструмент в руках творческого человека.

В этом качестве он не опасен, не имеет ограничений при решении как творческих задач, так и бытовых вопросов.

ИИ в этом качестве представляется мне как уникальный инструмент преобразования комбинаций пикселей, звуков, запахов и так далее одно в другое.

Если за такими преобразованиями стоят только творческие задачи, то можно увидеть очень большой простор для использования искусственного интеллекта.

Никаких перспектив использование ИИ не имеет в аналоговой среде. Он никогда не научится достоверно распознавать образы, управлять автомобилем, принимать судебские решения в спорте либо где бы то ни было.

Это связано с тем, что компьютер решает только конечномерные задачи, а реальная жизнь многогранна, непрерывна и неповторима.

Опасно применять ИИ в образовании, медицине, юриспруденции, и здесь нужны очень жесткие внутригосударственные ограничения.

Наконец, недопустимо использование ИИ при принятии решений, когда речь идет об опасных производствах, критических инфраструктурах, системах безопасности и так далее.

Говоря об использовании нейросетевого и генеративного искусственного интеллекта, необходимо отметить следующее:

- «хорошие», то есть обученные на большой выборке данных, модели ИИ – это всегда системный риск;
- во всех случаях ИИ учится на малой части смыслов и суждений, выработанных человечеством, имеет узкий кругозор в узкой предметной области, ангажирован;
- применение ИИ ограничивает теорема Гёделя;
- до сих пор не доказан эффект от внедрения ИИ: проведенные исследования показывают, что он не сокращает время труда специалистов;
- ИИ опасен в той же мере, что и любой закрытый код;

– результаты работы ИИ не позволяют отличить ложь от правды, так как результат является не смысловым, а статистическим;

– «функция ошибок» ИИ много выше обычных инженерных стандартов.

По мнению многих исследователей, для эффективной замены человека на ИИ в некоторых процессах деятельности требуется замена нейросетевых архитектур на какие-либо другие.

Чтобы отмахнуться от этих вопросов, адепты ИИ пиарят результаты работы ИИ и не дают оценок полезностям и опасностям этих результатов.

Из сказанного можно сделать две группы выводов.

Во-первых, по состоянию на сегодня:

1. Цифровая трансформация достигла такого уровня развития, на котором стал возможен сбор очень больших объемов статистических данных о протекании естественных, природных процессов в реальном масштабе времени.

2. Перспектива сбора все больших объемов статистики породила гипотезу об осуществимости управления реальной деятельностью людей при помощи статистических моделей в глобальном масштабе. К моделям, реализующим эту гипотезу, пока что отнесли:

– защиту озонового слоя Земли путем отказа от фреона;

– предсказание климатических изменений на основе статистики, собираемой с 1950 года;

– обещание создать искусственный (цифровой) интеллект, равный человеческому.

3. Соответствующие гипотезы агрессивно продвигаются в информационном пространстве, в том числе и при помощи нарочито создаваемой шумихи вокруг недостоверных и преувеличенных фактов о свойствах ИИ.

4. Важным элементом хайпа стало вручение Нобелевских премий в области физики за исследования озонового слоя, климата и искусственных нейронных сетей, хотя реальное применение соответствующих моделей не находит подтверждения на практике.

5. Главным физическим ограничителем распространения ИИ в виде цифровых статистических моделей является уровень энерго- и водопотребления на этапах обучения и эксплуатации моделей. ИИ требует в 10-15 раз больше энергии, чем традиционные алгоритмы вычислений, и усиливает реальный антагонизм между технологиями и природой. ИИ – враг устойчивого развития.

6. Бенефициарами шумихи вокруг искусственного интеллекта предполагают стать электронная промышленность США и военные (идеологические) агентства западных стран. Последние намерены получить в свое распоряжение более эффективные инструменты воздействия на общественное сознание, чем Голливуд, мессенджеры и социальные сети.

Во-вторых, необходимо ответить на вопрос: «Что делать?»

Как мне кажется, в России технологии ИИ следует развивать на основе принципов, сформулированных в Национальной стратегии развития искусственного интеллекта на период до 2030 года (Указ Президента РФ от 10.10.2019 г. № 490, содержание которого, кстати, меняется) и ряде других документов, принятых Федеральным собранием и Правительством РФ.

Для достижения целей национальной стратегии необходимо решить целый комплекс проблем:

– развитие информационно-коммуникационной инфраструктуры (сети, системы хранения и обработки цифровых данных);

- совершенствование сквозных информационных технологий (интернет вещей, большие данные, блокчейн, геоинформационные технологии);
- развитие Единой энергетической системы страны;
- разработка новых архитектур генеративных систем ИИ, лишенных недостатков нейросетевой архитектуры (непредсказуемость, высокое энергопотребление, системные риски применения);
- создание отечественного аппаратного, программного и информационного обеспечения систем ИИ;
- защита государства, бизнеса и личности от системных опасностей ИИ;
- выработка общепризнанного понятийного аппарата в этой области знания и так далее.

Большое спасибо за внимание.

Тосунян Г.А.: Спасибо.

Оба доклада, конечно, очень интересные.

Перейдем к вопросам.

Олег Анатольевич Ефремов, пожалуйста.

к. филос. н. ЕФРЕМОВ О.А. – д. т. н. КОЛБАНЁВ М.О.

ЕФРЕМОВ О.А.

к. филос. н., доцент кафедры социальной философии и философии истории философского факультета МГУ им. М.В. Ломоносова, ведущий эксперт Института фундаментальных проблем социогуманитарных наук Национального исследовательского ядерного университета «МИФИ»

Ефремов О.А.: Огромное спасибо Михаилу Олеговичу за прекрасный доклад.

Такая интересная ситуация складывается.

Когда об искусственном интеллекте говорят гуманитарии, это всегда что-то грандиозное – либо открывающее блестящие перспективы, либо, наоборот, ведущее нас к апокалипсису.

А вот технари часто улыбаются и в частных беседах говорят: на самом деле это всего лишь большой компьютер.

Михаил Олегович, у меня к Вам вопрос как к доктору технических наук.

В принципе, искусственный интеллект – это имитация некоторых интеллектуальных способностей человека.

Скажите, пожалуйста, есть ли какие-то человеческие способности, которые все-таки искусственный интеллект никогда не сможет удачно симитировать?

Есть ли что-то, доступное только естественному человеческому интеллекту, но принципиально недоступное искусственному?

Спасибо.

Колбанёв М.О.: Спасибо за вопрос, Олег Анатольевич.

Когда Вы его задавали, я подумал, что Вы так закончите: «А есть ли у искусственного интеллекта такие

возможности, которые могут использоваться для замены человека?»

А Вы инверсно это сформулировали.

Думаю, что сравнивать ИИ с человеком – это уже хайп.

Сопоставление человека и машины совершенно недопустимо.

Вы вспомнили специалистов-гуманитариев.

Когда я читаю их работы, в которых утверждается, что никогда искусственный интеллект не будет обладать сознанием, я про себя примерно так думаю.

Никто же не пишет о том, что табуретка никогда не будет обладать сознанием.

Все и так это понимают.

То же самое я думаю про искусственный интеллект.

Это просто большой калькулятор.

И перед инженерами, математиками, которые придумывают специальные сложные алгоритмы, надо снять шляпу, так как они делают этот калькулятор более понятным людям.

По крайней мере, то, что делается сейчас в этой области, позволяет создать, как мне кажется, мощный научно-технический задел на будущее, когда, может быть, и вычислитель станет еще мощнее, и, может быть, энергия не будет стоить так дорого, и, может быть, изменится что-то в человеческом обществе, в том смысле, как об этом Вавилов говорил.

Наступит какая-то эпоха разума.

Если все это возможно, тогда те заделы, которые сегодня мы получаем в этой области благодаря математикам и инженерам, там смогут использоваться.

Тосунян Г.А.: Спасибо. Пожалуйста, Александр Викторович Халявкин.

к. б. н. **ХАЛЯВКИН А.В.** – д. т. н. **КОЛБАНЁВ М.О.**

ХАЛЯВКИН А.В.

к. б. н., старший научный сотрудник

Института биохимической физики им. Н.М. Эмануэля РАН

Халявкин А.В.: Во-первых, большое спасибо Михаилу Олеговичу за такой познавательный во всех смыслах доклад.

У меня вопрос об информационном шуме, в том числе о намеренном вбросе ложной информации.

Всегда ли это несет в себе отрицательную коннотацию?

Потому что есть многочисленные примеры обратного. Спасибо.

Колбанёв М.О.: Да. Для кого-то разведчик, для кого-то шпион...

Вопрос в этом смысле?

Халявкин А.: Примерно. Во многих областях так.

Колбанёв М.О.: Искусственный интеллект, когда про него говорят в узком смысле, как о генеративном ИИ, – это и есть инструмент для вброса шума.

Во многом.

Причем достаточно иметь из тех рекомендаций, которые он дает, 5-10% ложной информации – это уже очень мощный шум.

Даже если в других 90% его рекомендации будут полезны.

Самое главное, что искусственный интеллект дает рекомендации, которые нельзя проверить.

Они статистические, они не несут смысла.

И поэтому, конечно, для специалиста, обладающего квалификацией в какой-то области знаний, он бессмыслен, потому что специалист сразу поймет, где там шум, а где что-то полезное.

А подавляющее большинство людей, которые не знают «акваторию» искусственного интеллекта, поверят, что в Черном море плавает кит.

Это может быть цель, ради которой внедряются зарубежные модели генеративного искусственного интеллекта в нашем обществе.

Для справедливости хочу сказать, что есть официальный запрет на использование генеративной модели OpenAI и некоторых других.

Может быть, я не все знаю.

Если коротко ответить на Ваш вопрос, это и есть инструмент для создания информационного шума.

Тосунян Г.А.: Спасибо.

Михаил Олегович, напомню, что Галина Вениаминовна запретила использовать словосочетание «ложная информация».

Сорина Г.В.: Простите, пожалуйста, но так как Вы сослались на меня, хотела бы сказать, что это, конечно, шутка, я ничего не могла запретить.

Но так как Вы меня пригласили выступить, я расскажу про это в своем докладе.

Спасибо.

Тосунян Г.А.: Хорошо.

В свое время, Галина Вениаминовна, я не в шутку рекомендовал, но не запретил использовать термина «социальная инженерия».

Я попросил моих коллег использовать выражение «социально-криминальная инженерия», потому что «социальная инженерия» очень нежно звучит для характеристики метода обмана, который используется для манипуляции людьми в преступных целях.

Мы, конечно же, рекомендуем, но никак не навязываем.

Слово Валерию Аркадьевичу Конявскому, пожалуйста.

д. т. н. КОНЯВСКИЙ В.А. – д. т. н. КОЛБАНЁВ М.О.

КОНЯВСКИЙ В.А.

д. т. н., заведующий кафедрой защиты информации
Московского физико-технического института

Конявский В.А.: Спасибо, Гарегин Ашотович. Представляюсь: я заведу кафедру защиты информации МФТИ.

Михаил Олегович, спасибо за прекрасный доклад. Мы по этому поводу еще немножко поспорим, когда будет возможность выступить.

Сейчас такой вопрос.

Я знаю, что такое финансовые пузыри, но не знаю, что такое «статистический пузырь» – термин, который Вы используете на втором своем слайде.

Что такое статистические пузыри?

Спасибо.

Колбанёв М.О.: Вы цитируете эмоциональную часть моего сообщения.

Это такая эмоциональная характеристика исследований, которые делают серьезные выводы на основе статистических данных, которые потом не подтверждаются и опровергаются другими статистическими данными.

Наверное, никто раньше такой термин не использовал.

Спасибо, что Вы на это обратили внимание.

Я об этом подумаю и, может быть, буду использовать в дальнейшем более понятный термин.

Конявский В.А.: Спасибо.

Тосунян Г.А.: Спасибо.

Коллега Смолин из Института прикладной математики имени Келдыша РАН.

СМОЛИН В.С.

научный сотрудник Института прикладной математики
им. М.В. Келдыша РАН

Смолин В.С.: Я тоже хочу поблагодарить докладчиков за очень интересный взгляд.

Конечно, хайп в области искусственного интеллекта присутствует, с этим невозможно спорить.

Особенно мне понравился пример про табуретку.

Очевидно, что у нее нет души, и это распространяется, соответственно, на все остальные предметы, произведенные человеком, в том числе на калькулятор.

Вы искренне считаете, что он отличается от человека тем, что у него нет души, поскольку Бог не дал?

И те вычисления, которые производит калькулятор, кардинально отличаются от тех вычислений, которые производит человек?

Потому что и то и другое – интеллектуальная деятельность.

В чем Вы видите принципиальное отличие тех результатов, которые получены вычислением в столбик, и калькулятором?

Колбанёв М.О.: Спасибо Вам за очень сложный вопрос.

Он касается на самом деле многих явлений, которые связаны с внедрением современных цифровых технологий.

Различия, наверное, есть.

Припоминаю, что вычисления в Древнем Риме выполняли рабы, потому что патриции не могли опуститься до такой простой работы.

Ну, теперь эти вычисления выполняют калькуляторы, заменив этих рабов.

Может быть, лучше привести другую аналогию.

Человек, когда он производит вычисления, не просто вычисляет.

За этой его деятельностью, связанной с вычислением, стоит еще что-то – то, что очень трудно описать, и то, чего не сможет сделать компьютер.

Напомню заявление одного из наших самых крупных бизнесменов, можно сказать, и государственного деятеля (фамилию называть не буду), который лет 5-7 назад обещал заменить 40 тысяч своих юристов на искусственный интеллект.

Вы знаете, никого не заменил. Потому что, оказывается, эти юристы помимо того, что заполняют стандартный бланк договора с клиентом, делают еще что-то, что не позволяет их заменить.

Отличие именно в этом.

Человек, выполняя вычислительные задачи, не просто вычисляет – он еще и реализует какие-то свои человеческие качества.

Этого, мне кажется, нельзя не заметить.

Я это понимаю, хотя мне это очень трудно объяснить.

И когда соединение устанавливает телефонистка, а не автоматическая телефонная станция, то в этом тоже есть что-то особенное. Хотя внешне видно, что они выполняют одну и ту же функцию.

Между прочим, о Шалапине многое известно только потому, что его слушали в Петербурге телефонистки, а он много пользовался телефоном.

Если коротко, любые вычисления, которые выполняются человеком, – это не только вычисления, это еще какие-то действия, которые поддерживают общий процесс деятельности.

Тосунян Г.А.: Спасибо, Михаил Олегович.

А если еще короче, я бы сказал так.

Вычислительные операции человек осуществляет с душой, а калькулятор – без души.

Пожалуйста, Виктор Дмитриевич Кривов.

д. э. н. КРИВОВ В.Д. – д. т. н. КОЛБАНЁВ М.О.

Кривов В.Д.: Спасибо большое, уважаемый Гарегин Ашотович.

И огромное спасибо, уважаемый Михаил Олегович, за такой достойный доклад.

У меня два вопроса.

Один вопрос не совсем по теме.

Может быть, Вы сталкивались с какими-то исследованиями...

Я общался с коллегами из МИРЭА. У них уже преподаются дисциплины, как использовать ИИ для решения разных задач и как разрабатывать ИИ, нейросети.

По их оценкам, наиболее успешно продвигаются в решении этих задач студенты, подверженные аутизму.

Это первый вопрос.

И второй вопрос связан с утечкой информации.

Любой, кто формулирует запрос нейросети, соответственно, оставляет свой след, и нейросети потом даже формируют профили своих пользователей для каких-то целей.

Спасибо.

Колбанёв М.О.: Спасибо большое. Сложные вопросы.

Я не так глубоко понимаю особенность аутизма.

Но предположу, что просто эти люди больше ориентируются на свои ощущения, чем на внешнее воздействие, которое так или иначе оказывается на человека в социальной среде.

Может быть, с этим связана более легкая, эффективная работа аутистов с искусственным интеллектом.

Они больше чувствуют себя и исходя из этого создают те компьютерные программы или системы, которыми они занимаются.

К сожалению, не могу более подробно об этом сказать. Исследований таких не знаю, я об этом не задумывался и никогда ничего об этом не читал.

Что касается того следа, который оставляет у себя в памяти компьютерная система о каждом своем пользователе, то это непреложный факт, в этом не нужно сомневаться.

И это касается не только генеративных систем, искусственного интеллекта, а вообще всех цифровых систем.

Чем больше система знает о человеке, тем более удобную услугу она может ему предложить.

Поэтому многие люди охотно сообщают о себе системе некоторые данные, личные или, может быть, даже интимные, для того чтобы в результате услуга была удобна и проста, чтобы она предоставлялась нажатием одной клавиши.

Но чем больше информации мы даем о себе или система узнает о нас, тем большие риски использования этого интеллекта, в том числе и риски для частной жизни человека.

Границу этих противоречий – удобство, с одной стороны, и безопасность, с другой стороны – сегодня каждый для себя должен провести сам.

Хотя, возможно, это должно стать и предметом регулирования на государственном, законодательном уровне.

Это большая проблема вообще всех цифровых систем. То, о чем Вы спросили, – это очень важно.

Кривов В.Д.: Спасибо большое.

Тосунян Г.А.: Спасибо.

Алина Сергеевна Зайкова, представьтесь, пожалуйста.

к. филос. н. ЗАЙКОВА А.С. – д. т. н. КОЛБАНЁВ М.О.

ЗАЙКОВА А.С.

к. филос. н., научный сотрудник Института философии и права
Сибирского отделения РАН

Зайкова А.С.: Добрый день!

Большое спасибо за приглашение на это интересное заседание. И большое спасибо за доклады.

У меня два коротких вопроса.

Первый связан с Вашей классификацией безопасных и опасных областей.

Мне кажется, что даже если это безопасная или опасная область, в любом случае ответственность должна лежать на операторах и агентах.

То есть редактор журнала смотрит и принимает решение, берет он это изображение для печати или не берет.

Врач смотрит результаты какого-то ИИ-исследования и в дальнейшем решает, как он их оценивает, и так далее.

Соответственно, именно человек ответствен за принятие решения. Разве это не так?

Это первый вопрос.

Второй вопрос относится к примеру с табуреткой.

Мы, действительно, не говорим, что у табуретки может быть сознание, даже не задаемся таким вопросом.

Но табуретка стоит на месте, она не разговаривает и никаких действий не совершает.

Но если мы, допустим, обратим внимание на роботов-пылесосов, то большинство людей, которые их используют, с ними разговаривают.

Если мы обратимся к опыту использования голосовых помощников, той же умной колонки Яндекс.Алиса, то, опять же, люди с ней общаются.

И это общение выходит за рамки простого получения информации.

В большинстве случаев люди совершенно не ожидают, что у голосового помощника или у робота-пылесоса есть душа.

Но тем не менее некое иное общение с роботом-пылесосом или голосовым помощником, отличное от общения с табуреткой, присутствует.

Не кажется ли Вам, что сама постановка вопроса, может ли быть сознание у искусственного интеллекта, связана именно с этим фактором социального взаимодействия?

Благодарю.

Тосунян Г.А.: Спасибо, Алина Сергеевна, за вопросы.

На второй, как мне кажется, ответить проще.

Когда люди начинают относиться к гаджетам или к пылесосам как к домашним питомцам, то это и есть результат хайпа, который создан вокруг искусственного интеллекта.

Это и есть тот самый информационный шум, который обманывает многих людей.

Пылесос никогда не будет котом, собачкой либо другом. А еще он не будет ни мужем, ни дедом, ни другом. В общем, как и табуретка. И не надо к ним относиться как к живым существам.

А первый вопрос...

Зайкова А.С.: Про то, что в любых случаях, и в опасных, и в безопасных областях, ответственность лежит на операторе.

Колбанёв М.О.: Спасибо.

Вообще, это вопрос очень сложный. Простых вопросов в этой области, наверное, и не бывает.

Как регулировать ответственность за использование искусственного интеллекта?

Самые известные примеры связаны с беспилотными автомобилями, которые нарушили правила...

И кто в этом виноват? Конструктор, программист, пешеход или кто-то еще?

Вы понимаете, здесь спектр ответов.

Пока их нет, но можно для каких-то широких предметных областей, может быть, принять законы, где ограничить уровень принятия решений, который разрешается машине.

И установить, что с какого-то уровня принятия решения необходимо обязать производителей привлекать людей для подтверждения этого решения.

В этом направлении сделан определенный шаг и в законе ЕС об искусственном интеллекте (AI Act), который я упоминал, и в указе президента США, еще Байдена, который очень похож на европейский.

В них дается такая оценка.

Опасность ИИ зависит от того, сколько вычислительных операций было реализовано для его обучения.

Если для обучения ИИ использовали «много» операций, то он опасен и должен получать специальное разрешение для внедрения в практическую деятельность.

В Европе число таких операций равно 2^{25} -й флоп (компьютерных операций с плавающей точкой), а в Штатах – 2^{26} -й степени.

И то и другое очень много. Самый быстрый суперкомпьютер на стандартном тесте выполняет 2^{18} таких операций в секунду.

Если разделить 2^{26} на 2^{18} , получим 2^8 (двести миллионов) секунд, а это больше шести лет обучения.

Насколько я знаю, суперкомпьютеры не используются для обучения ИИ, но приведенный пример дает представление о дороговизне и сложности процессов обучения ИИ.

Если обучение требует большого числа компьютерных операций, тогда государству нужно считать, что эта система несет в себе системный риск. И нужно проверить, насколько он все-таки опасен и можно ли его вообще использовать.

О том, как практически организовать проверку «опасности» ИИ, в документах не говорится.

Если операций немного, тогда, наверное, опасности меньше, однако и толку от использования такого ИИ много меньше.

Но вопрос, Алина Сергеевна, который Вы задаете, неоднозначный, он очень сложный.

Это все предстоит изучать и разработчикам искусственного интеллекта, и инженерам, и гуманитарным специалистам, и всему научному сообществу.

Спасибо еще раз за вопрос.

Зайкова А.С.: Большое спасибо.

Тосунян Г.А.: Спасибо.

Академик Барановский, пожалуйста.

акад. БАРАНОВСКИЙ В.Г. – д. т. н. КОЛБАНЁВ М.О.

БАРАНОВСКИЙ В.Г.

акад. РАН, акад. Академии военных наук, д. и. н., научный
руководитель Центра ситуационного анализа Национального
исследовательского института мировой экономики
и международных отношений им. Е.М. Примакова РАН

Барановский В.Г.: Спасибо. Хочу поблагодарить
Гарегина Ашотовича за организацию этого семинара и
Михаила Олеговича – за великолепный доклад, который я
слушал с огромным вниманием.

Я занимаюсь международными делами и могу засви-
детельствовать.

В области международных отношений понятие ис-
кусственного интеллекта используется все шире, очень ак-
тивно. И мы сталкиваемся примерно с такими же пробле-
мами, как те, о которых говорил Михаил Олегович.

Мне кажется, что в докладе Михаила Олеговича обо-
значены очень важные реперные точки, и я хотел бы его
активно поддержать.

Шум вокруг ИИ часто инициируется людьми, кото-
рые искренне стремятся к расширению понимания дей-
ствительности, к выявлению каких-то новых закономер-
ностей.

Но нередко есть и другие мотивы, которые дали
нашему докладчику основание использовать термин
«хайп».

По-моему, очень интересный и очень точный термин.

У меня есть предложение.

Учитывая, что это связано и с деньгами, и с преми-
ями, и с имиджем, и с немалыми материальными ресур-
сами (включая государственное финансирование), не

хотите ли Вы в качестве дополнительного использовать термин «панама»?

Это мое предложение для Ваших рассуждений на эту тему.

Спасибо.

Колбанёв М.О.: А почему нет? Может быть.

Я могу с этим согласиться и подумать об этом.

Спасибо за предложение.

Тосунян Г.А.: Спасибо.

Коллеги, есть еще вопрос Льва Московкина. Пожалуйста, Лев.

МОСКОВКИН Л.И.

специальный корреспондент газеты «Московская правда»

Московкин Л.И.: Спасибо большое.

Для чего озоновые дыры, я знаю.

Для чего глобальное потепление, тоже.

Искусственный интеллект, его «шумовое» оформление...

Правильно ли мое впечатление, что для размывания государства?

Генеративный – это и есть самопрограммируемый.

Аутизм – это очень редкая болезнь.

Когда формируется мозг, образуется огромное количество синапсов. Потом они начинают удаляться, остаются только те, которые работают.

При аутизме этого не происходит. Ребенок сам в себе живет.

Но спасибо большое.

Ваш доклад мне особенно понравился.

Спасибо.

Тосунян Г.А.: Спасибо.

Переходим к обсуждению, поскольку вопросы уже задали обоим докладчикам. А докладчиков я прошу в процессе выступления наших коллег посмотреть вопросы в чате, чтобы ответить на них в ходе дискуссии или в заключительном слове.

Коллеги, давайте послушаем Валерия Аркадьевича Конявского, заведующего кафедрой в МФТИ. Пожалуйста, Валерий Аркадьевич, Вам слово.

д. т. н. КОНЯВСКИЙ В.А. – д. п. н. СЕРГЕЕВ С.Ф.

Конявский В.А.: Спасибо, Гарегин Ашотович.

Мне кажется, что Михаил Олегович своим докладом достиг основной цели, которую ставил перед собой, – обострить дискуссию.

Дискуссия обострена, что позволило увидеть проблему с новой точки зрения.

Взгляд, который мы все увидели, безусловно, новый и очень самобытный.

У нас с Михаилом Олеговичем абсолютно точно совпадают исходные данные.

Я с недоверием отношусь к статистическим моделям, довольно плохо отношусь к нынешнему уровню развития технологий искусственного интеллекта.

И на этом почти все совпадения у нас заканчиваются.

Кстати, очень здорово получить заранее презентацию докладчика, чтобы с ней ознакомиться.

Большинство тезисов презентации и, соответственно, доклада вызывают у меня некоторый протест.

Об этом я и хотел поговорить.

Если бы мне нужно было дать название моему короткому выступлению, я назвал бы его примерно так: «Стакан наполовину полон».

В этом отличие от того, что мы слышали от Михаила Олеговича: якобы стакан пустой, и ничего, кроме дыма и хайпа, там как бы нет.

В моем представлении стакан наполовину полон.

Я серьезно думал о том, почему некоторые тезисы вызвали у меня такой протест.

И понял.

Потому что предметом этого доклада стали бытовые представления о современном уровне развития ИИ.

Бытовые, а вовсе не те, которые на самом деле функционируют в науке, в той, в которой нам приходится «крутиться» в части обеспечения безопасности и доверенности систем ИИ.

А если мы говорим о безопасности и немного понимаем технологии, то некоторые тезисы доклада, на мой взгляд, могли бы быть уточнены.

Во-первых, дезинформация – это угроза.

Это, безусловно, так. Я двумя руками поддерживаю этот тезис. Это совершенно точно и крайне важно.

Как-то пару лет назад мы написали книгу, посвященную мультиагентным системам.

В ней показано аналитически, что если 10, 12, 15% дезинформации «впрыскивается» в мультиагентную систему, то она очень быстро деградирует и потом погибает.

Она умирает просто потому, что есть дезинформация.

Вывод сам по себе очень непростой и даже опасный, потому что в нашем мире довольно много дезинформации.

К чему это приведет?

Ну, хотелось бы жить подольше.

Тосунян Г.А.: Погибает система? Не информация погибает?

Коняевский В.А.: Нет, нет, система. Что касается информации, знаний и так далее...

Сергеев С.Ф.: А какая система погибает?

Коняевский В.А.: Мультиагентная система.

Сергеев С.Ф.: Что значит погибает?

Коняевский В.А.: Перестает функционировать. Она не выполняет свою функцию.

У мультиагентной системы есть функции.

Вот, например, клин журавлей.

Журавли летят клином.

Не потому, что так легче преодолевать сопротивление ветра, а потому, что у журавлей глаза устроены по-разному: некоторые стремятся видеть соседа слева, другие – справа, третьи – небо вверху и землю внизу.

Размещаясь в пространстве по этому принципу, они выстраиваются клином и летят горизонтально.

Если нарушить этот принцип, клина не будет.

И функция не будет реализована.

Другое дело, что галлюцинации в нынешних «болталках», в LLM – больших языковых моделях, возникают очень быстро, очень часто.

Они, конечно, шум, они, конечно, мешают.

И об этом мы поговорим.

Хочу вернуться к статистике.

Любая модель не точна. Любая модель – это упрощение.

Статистика не исключение.

Это модели, которые не точны, они могут давать неверный результат.

Но без моделей нет прогнозов, без прогнозов нет управления, а без управления нет эффективности.

Остается только что?

Энтропия⁸.

Вот тут система и начинает погибать.

⁸ Мера беспорядка, хаоса, степень неопределенности.

Поэтому я не стал бы отрицать статистические модели как таковые. Я стал бы думать о том, что не так, что можно улучшить в статистических моделях.

Статистические модели могут улучшаться с позиций доверенности, мы работаем над этим, но это тема отдельного обсуждения.

Теперь – что касается информационного шума. Он не равен дезинформации.

Дезинформация – целенаправленное искажение.

Информационный шум – это не обязательно дезинформация.

Опасно и то и другое: мешают и шум, и дезинформация.

Очень важна представленная Михаилом Олеговичем классификация базовых объектов цифровой трансформации.

В моем представлении она чуть иная, и я думаю, что есть смысл об этом поговорить отдельно.

Мне кажется, что раньше суть информационных систем касалась документоцентричности, то есть базовым элементом считался документ.

А сейчас, по прошествии лозунгов «цифровизации», «цифровой экономики», «экономики данных», правильный термин – это датацентричность, то есть теперь известно, что основным объектом цифровизации являются не документы, а данные.

Это принципиально.

Спасибо Михаилу Олеговичу за очень остроумные наблюдения. Например, обучение проходит сегодня на базе того, что называется Data Lake.

То есть мы сначала собираем огромное количество данных, сливаем их в некоторое «озеро данных» (Data Lake), а потом прикладываем огромные усилия, огромную энергетику, огромные вычислительные мощности, для того чтобы из этой взвешенной смеси вытащить те данные, которые нам почему-то нужны.

Это абсолютная бессмыслица, это абсолютно неправильный подход к искусственному интеллекту, просто потому что сначала данные смешиваются с нарушением семантических связей, а потом с огромными усилиями и затратами приходится пытаться их разделить.

И это замечание я считаю очень полезным и нужным.

Нет точного определения искусственного интеллекта. Действительно, нет. Общепринятого – тем более нет. Но это не значит, что его не может быть.

Могу сослаться только на известный афоризм, приписываемый Фаине Раневской, по поводу того, что как бы некоторая сущность есть, а слова нет, или слово есть, а сущности нет.

Здесь я не стал называть ту сущность, о которой говорила Раневская, но я думаю, что это более-менее понятно.

Теперь – об архитектуре искусственного интеллекта.

Ведь мы рассматриваем систему с точки зрения ее влияния на мир.

И с этой точки зрения более точной классификацией мне представляется следующая: необходимо разделить систему ИИ на информационное, киберфизическое и автоматическое взаимодействие.

Понятно, что эти виды взаимодействия совершенно различны.

Информационное взаимодействие на физический мир не влияет, влияет только на восприятие человека.

Киберфизическое взаимодействие происходит с участием человека, автоматическое – совсем без участия человека.

На такой классификации гораздо легче строятся механизмы взаимодействия с системами искусственного интеллекта.

Искусственный интеллект не равен генеративному искусственному интеллекту.

Искусственный интеллект намного шире, чем просто генеративный, и сводить ИИ к большим языковым моделям, мне кажется, тоже не стоит.

Нынешняя схема реализации больших языковых моделей, то есть LLM, неизбежно приводит к галлюцинациям. Это так, потому что применяемые механизмы обучения без учителя, на материалах «озера данных», аналогичны попытке поговорить в трамвае.

Там едут разные люди, у всех разные интересы, и все говорят о своем, а слышим мы то, что хотим слышать.

Для того чтобы генеративная модель галлюцинировала поменьше, к ним добавляются механизмы управления контентом, так называемый RAG.

Но механизм управления контентом – это примерно «а давайте поговорим вот об этом».

Но говорят все равно разные люди.

Чтобы исчезли галлюцинации, необходимо в начале системы поставить то, что называется системой управления знаниями.

То есть мы говорим об этом в таком-то контексте, хотим выявить знания вот об этом, цель нашего разговора вот такая.

Тогда галлюцинации исчезают.

Это, я думаю, механизм развития больших языковых моделей, который ожидает своей реализации в ближайшем будущем.

И почему вообще такая ошибка произошла?

Эта ошибка, на мой взгляд, связана с тем, что я назвал бы ошибкой великого Тьюринга.

Ведь Тьюринг придумал вычислительную машину, которая может решить любую задачу.

Но если она решает любую задачу, то она решает и вредоносную задачу.

Это архитектура, которую нельзя избавить от ошибок в силу архитектурных особенностей.

Мы прекрасно знаем, что целое – это единство формы и содержания.

Цифра – это только форма.

Вот отсюда возникают галлюцинации.

Содержание – это семантика.

Если мы добавим семантику к тому, что сегодня получаем на вычислительных машинах, тогда с калькулятором нас будет сравнивать довольно трудно.

5 плюс 6 – это всегда 11?

Вряд ли, потому что если это 5 помидоров и 6 яиц, скорее всего, это будет яичница, если на сковороде. А если в холодильнике, то так и останется 5 плюс 6.

То есть форма без содержания практически не работает.

В этом недостаток развития современной вычислительной техники, но я думаю, что он может быть преодолен.

Перейду непосредственно к выводам.

С первым выводом Михаила Олеговича я согласен полностью, тут не может быть даже никаких сомнений.

Со вторым выводом я уже никак не могу согласиться, потому что климатические изменения к деятельности людей имеют очень опосредованное отношение, поэтому с климатом не очень точно получилось.

Что касается вывода номер три, не могу его принять, но достаточного времени на дискуссию сейчас нет. Если нужно, мы это с Михаилом Олеговичем отдельно обсудим, потому что подход очень интересный, и обсуждение было бы полезным.

Четвертый. К сожалению, в моем представлении вывод о том, что элементом хайпа стало вручение Нобелевских премий, немного похож на ту дезинформацию, которую Михаил Олегович так рьяно критиковал в своем докладе.

Ведь Нобелевская премия вручена не за искусственный интеллект по физике.

Она вручена за создание на новых физических принципах ассоциативной памяти, которая может быть использована, в том числе в вычислениях по искусственному интеллекту.

Это небольшая, на мой взгляд, тонкость, которую имело бы смысл указать.

«Искусственный интеллект – враг устойчивого развития».

Это пятый вывод.

Здесь должен напомнить только одно: не так давно на нашей памяти таким врагом была кибернетика.

С шестым пунктом, конечно же, нужно согласиться по той простой причине, что, на мой взгляд, нет ничего плохого в том, что будет развиваться электроника.

Нам нужно развивать ее, а не ждать, пока за нас это сделают другие.

Итак, развитие ИИ связано с внедрением семантики в вычисления, которые сегодня производятся только по форме, без учета содержания.

Это влияет в свою очередь на развитие статистических методов.

Поэтому методы в статистике сегодня активно развивают так называемые связанные данные, что само по себе представляет внедрение семантики и семантической интероперабельности в статистические системы.

Здесь не ИИ плохой, потому что статистические модели плохие, а изучение ИИ позволяет улучшить статистические модели пониманием того, что без статистики не обойтись.

Спасибо огромное.

Тосунян Г.А.: Спасибо.

Пожалуйста, Олег Анатольевич Ефремов.

к. филос. н. ЕФРЕМОВ О.А. – акад. ТОСУНЯН Г.А.

Ефремов О.А.: Извините, я в первый раз не представился. Ефремов Олег Анатольевич, философский факультет МГУ им. М.В. Ломоносова.

Несколько замечаний по поводу очень интересных докладов, которые сегодня были, и обсуждения, которое потом последовало.

Конечно, много хайпа вокруг искусственного интеллекта, с этим спорить невозможно.

Но все же за этим хайпом что-то стоит. Как говорится, нет дыма без огня, и все-таки огонек здесь тоже присутствует.

Первое.

Как философ, я все-таки вынужден констатировать, что процесс цифровизации, частью которого является появление искусственного интеллекта, по сути дела, создает новое царство бытия.

До сих пор этих царств было несколько: неживая природа, живая природа, общество, человек.

И всякий раз мы конкретизировали философские категории применительно к каждому из этих царств.

И даже «полномочия» кафедр философского факультета делятся соответствующим образом.

Сегодня мы вынуждены осуществлять новую конкретизацию философских категорий применительно к той самой особой цифровой среде, которая сегодня уже несводима ни к чему из того, что есть.

Действительно, сначала мы говорили всего лишь о технологиях в рамках общества.

Потом появились теории постиндустриальных информационных обществ, где технологии уже создают тип обществ.

А теперь мы уже говорим о возникновении совершенно новой реальности, пусть выросшей из социальной, человеческой, но отличной от них, выстраивающей с ними очень сложные отношения.

Говорят даже о технологической сингулярности, когда эта реальность полностью оторвется от человеческой и социальной и обретет совершенно самостоятельную жизнь.

Поэтому значение цифровых процессов и искусственного интеллекта как их части преуменьшать нельзя.

Второе.

Искусственный интеллект.

Да, определений, конечно, много.

Но я все-таки полагаю, что это прежде всего имитации неких возможностей, способностей человеческой психики, которые осуществляются комплексами, созданными на основании цифровых технологий, и которые, в принципе, могут выступать аналогами некоторых функций человеческой психики.

Но что не в состоянии, на мой взгляд, имитировать искусственный интеллект с достаточной степенью успешности?

Михаил Олегович говорил, что ИИ способен к творчеству.

Как раз в этом я очень сомневаюсь, потому что творчество – это не просто перекомбинация того, что было.

Творчество – это создание чего-то принципиально нового.

Подлинное творчество – это прорыв.

Великое научное открытие, наверное, не может быть скомбинировано просто на основании каких-то суммированных прошлых знаний.

Вряд ли на это способен ИИ – здесь нужна та самая человеческая гениальность, которой у него нет.

И второе – это осознанный выбор, с которым связана свобода воли.

Да, искусственный интеллект, видимо, может выбирать между вариантами.

Но этот выбор всегда должен быть каким-то образом ценностно ориентирован.

Видимо, это пока – или, может быть, в принципе – тоже способность только человеческая.

Кстати, по поводу управления автомобилем. Я очень сомневаюсь, что ИИ не может этого делать.

Я пять минут ехал на машине, управляемой автопилотом, по горному серпантину ночью, причем дорога была загруженной.

Нет, я не такой сумасшедший, за рулем сидел другой человек, это была его машина, он включил автопилот, и пять минут мы благополучно ехали.

Как видите, я жив и могу сегодня выступить перед вами.

Так что, видимо, искусственный интеллект автомобилем управлять сможет и судить сможет, а вот к творчеству и осознанному выбору, который составляет основу свободы воли, вряд ли способен.

И наконец, третье, очень важное.

Когда мы говорим о возможностях ИИ и где-то пакуем перед ним, полагая, что он сильнее нас, мы констатируем свою слабость, мы констатируем собственную деградацию.

К сожалению, разрушение когнитивных способностей человека в XX – начале XXI века подготовило нас к тому, что ИИ становится сильнее нас.

Мы, по сути дела, убираем у себя возможности, которые отличают нас от ИИ, и потому ему проигрываем.

Приведу простой пример.

По какому критерию мы оцениваем сегодня качество выпускных работ наших студентов?

Сколько концепций они пересказали, на кого они сослались, что они задействовали.

Разве ИИ не может осуществить такую компиляцию?

Вполне.

Мы не ценим самостоятельность, не призываем к собственному мышлению, собственному творчеству.

Точнее, мы можем это говорить, но работы-то по факту оцениваем по-другому.

Я знаю даже докторские диссертации, которые написаны именно так. Там ни одной своей мысли, ничего нового, только добросовестный пересказ того, что говорили другие.

Да, это большая работа, но эту работу вполне в состоянии выполнить ИИ.

И это только один из примеров.

Искусственный интеллект побеждает нас там, где мы деградируем.

Он побеждает нас там, где мы не развиваем свои собственные человеческие возможности, отличные от него.

И искусственный интеллект может быть не силой, опасной для нас, он может быть не тем, кто нас победит, а может и должен быть тем средством, которое мы используем.

Но для этого мы интеллектуально должны быть сильнее его.

Искусственный интеллект, созданный нами, — это тоже проявление нашего могущества.

Но это сегодня и вызов по отношению к нам.

Человек должен быть сильнее, умнее и ответственнее тех технологий, которые он создает.

Спасибо.

Тосунян Г.А.: Спасибо, Олег Анатольевич.

Но Ваш пример в качестве доказательства, что Ваше присутствие в целости и невредимости на нашем «Рабочем завтраке» — аргумент в пользу искусственного интеллекта, убедителен.

Здесь правильнее благодарить Господа Бога за то, что Вас пощадил.

Ефремов О.А.: Гарегин Ашотович, я никогда не полагаюсь на ИИ за рулем, к тому же в моей машине его нет.

Просто в данном случае мы можем говорить о том, что мой пример опровергает утверждение Михаила Олеговича о том, что ИИ не может управлять машиной.

Видимо, все-таки совершенная система управления может появиться.

Тосунян Г.А.: Может быть, он всего-навсего говорил о том, что там другие риски, и непонятно, чья ответственность? Но сейчас не будем уходить в детали...

Виктор Дмитриевич Кривов, пожалуйста, Вам слово.

КРИВОВ В.Д.

Д. Э. Н.

Спасибо, Гарегин Ашотович.

Уважаемые коллеги, на экономическом факультете МГУ им. М.В. Ломоносова уже третий год действует решение Ученого совета разрешить студентам использовать нейросети при подготовке своих курсовых, лабораторных, выпускных работ, в том числе и магистерских диссертаций.

Но вот терминология мне кажется все-таки важной.

Приведу такой пример.

Цифровизация экономики, вообще цифровизация, при этом опора идет, видимо, на арабский язык.

В латинском языке есть схожее слово, но оно заимствовано из арабского и обозначает ноль, однако «нулефикация экономики», конечно, смешно звучит.

Тем не менее *Artificial Intelligence* переводится на русский язык, очевидно, все-таки неправильно.

Тем более что слово *intelligence* в английском имеет много значений.

Например, в Высшей школе экономики есть магистерская программа «Analytic BI», то есть *Business Intelligence*. В этом случае словосочетание переводится как «деловая разведка».

Есть термин «генеративный», к которому мы постепенно привыкаем... В русском языке давно было слово «генерирующий». А «генеративный» употребляется сейчас уже меньше, это слово ассоциируется со словом «дегенеративный».

Я прошу прощения.

Когда идут дискуссии по этому поводу, мне всегда вспоминается известное произведение – «Золотой теленок».

Там было государственное учреждение «Геркулес», где был директор Полыхаев.

Однажды Бендер засел у Полыхаева с утра и на весь день, а секретарша тем временем сильно волновалась, потому что было много очень важных, срочных бумаг.

Но в запасе у нее были резиновые штампы.

И формулировки этих резолюций на штампах постоянно пополнялись, развивались, уточнялись, так что можно сказать, что секретарша Полыхаева в системе в совокупности с этими штампами вполне заменяла нынешний искусственный интеллект.

И авторы «Золотого тельца» сделали вывод, что «резиновый» Полыхаев ничуть не хуже живого.

И еще хочется немного поговорить об опасностях, в частности о том, что называется утечкой информации.

Наверное, все мы помним шахматные компьютеры.

Тогда каждая команда разработчиков обучала свой компьютер от соревнования к соревнованию, от поединка к поединку.

Но тогда обучение конкретного компьютера шло в замкнутом пространстве.

А теперь получается, что у нас нейросети ведут сеанс одновременной игры, и при этом вполне можно допустить, что группа мошенников и группа полицейских, которые пытаются выявить и обезвредить этих мошенников, могут обращаться к одной и той же нейросети.

Она будет обучаться и делиться полученными знаниями то с одной группой, то с другой в режиме реального времени.

Это, пожалуй, то, что я хотел сказать в порядке реплики. Спасибо.

СОРИНА Г.В.

д. филос. н., проф.

Спасибо большое, Гарегин Ашотович.

У меня очень короткая реплика.

В рамках обсуждения звучала мысль о том, чтобы запретить искусственному интеллекту принимать решения.

Я хотела сказать, что это в принципе невозможно.

Сама теория принятия решения возникает исходно в 1940-е годы XX века как теория игр, которую начинают в такой модели, как принятие решения, разрабатывать Фон Нейман и Оскар Моргенштерн.

И на этой базе появляется так называемая нормативная модель принятия решения.

Эта нормативная модель принятия решения опирается в том числе на искусственный интеллект.

Другое дело, что есть другая модель принятия решения – дескриптивная модель принятия решения, которая, с моей точки зрения, как раз и должна контролировать то, что происходит в рамках нормативной теории принятия решения.

Дескриптивная модель принятия решений не нуждается в математическом моделировании.

В дескриптивной модели описывается реальное положение дел в рамках обсуждаемых проблем в повседневности, науке, политике, международных отношениях.

В этом описании есть альтернативы.

Задача заключается в том, чтобы представить аналитику, позволяющую обосновать выбор той или иной альтернативы.

Описательный характер такой модели не означает того, что она не является научной.

Научные дискуссии как раз и проходят в рамках дескриптивных моделей обоснования и принятия решений.

И самая последняя реплика.

Идеи дескриптивной модели принятия решения фактически впервые появляются в книге «Никомахова этика» Аристотеля.

Его идеи в этой области, как и во многих других областях, думаю, остаются актуальными.

Спасибо.

Тосунян Г.А.: Спасибо.

Дмитрий Николаевич Кавтарадзе, биологический факультет МГУ им. М.В. Ломоносова. Пожалуйста, Дмитрий Николаевич.

КАВТАРАДЗЕ Д.Н.

д. б. н., ведущий научный сотрудник кафедры общей биологии
и гидробиологии биологического факультета
МГУ им. М.В. Ломоносова

Хотел бы поблагодарить коллегу, философа А.Ю. Алексеева, за его неотступное желание привлекать естественников к обсуждению философских проблем.

На мой взгляд, справедливы все высказанные комплименты, и их можно дополнить.

О чем мы немного забыли?

О задаче, которая возложена на наш круг, – заботиться хоть как-то о культуре масс или массовой культуре популяции.

Это различные вещи.

Светская картина мира у современного студента – философы меня, может быть, простят – отсутствует.

Оканчивая в советское время МГУ, мы имели мировоззрение и знали примерно, кто мы и где, а также как нам себя вести.

Сейчас этого нет, эту функцию продолжают выполнять только конфессии.

Никто об этом сегодня не сказал.

Мы пришли к тому, что проблемы состояния окружающей среды, биосферы – контекста эволюции – многим неизвестны.

Мало того, многие сограждане вообще забыли, что мы живые, и биологи часто слышат: ну, что нам об этих инфузориях или муравьях каких-то талдычат, а мы-то здесь при чем?

Несколько поколений не успело узнать, да и общество в целом подзабыло, что и мы – живые.

В школах страны не было мировоззренческих предметов: астрономии – около 15 лет, экологии – около 30 лет.

Поэтому проблемы и сценарии «пределов роста» Римского клуба, модель «ядерной зимы» Н.Н. Моисеева оказались отодвинутыми, а масштаб событий в мире не освоен.

Международные организации и отдельные страны не успели⁹ выбрать и найти инструменты снижения рисков развития, обеспечения «устойчивости».

Нам остаются вера и поиск – таково состояние мировоззрения.

Говоря об ИИ, никто не упомянул, к сожалению, о том, что полвека существуют способ системного мышления и такой язык, как системная динамика¹⁰.

Существуют также «нормальные аварии»¹¹, поскольку всякая сверхсложная система при отказе одного или двух элементов проявляет собственное поведение, которое научно не удастся предвидеть.

Очень важно, что сегодня многие обсуждали смысловые потери, или редукцию смысла в управлении.

Посмотрим на «пирамиды управления», которые мы создали, включая нашу «Истину», и признаем, что в системе имеет место редукция смысла.

Чем больше уровней, тем больше посредников в «редактуре» указаний сверху и сообщений снизу.

С детства знакомая многим игра «Испорченный телефон» – тому доказательство.

⁹ О принципах профессора Денниса Медоуза см.: Кавтарадзе Д.Н., Ямова Е.А. От устойчивого к сопряженному развитию // Природа. 2024. № 7.

¹⁰ Донелла Медоуз. Азбука системного мышления. М., 2010.

¹¹ Chales Perrow. The Normal accidents: Living with High-Risk Technologies, 1984, 1999.

д. б. н. Кавтарадзе Д.Н.

Надеюсь, что семинары продлятся, и проблемы игнорирования нами эволюции биосферы и нас как ее продукта займут в них основополагающее место.

Спасибо большое.

Тосунян Г.А.: Спасибо.

Пожалуйста, Смолин Владимир Сергеевич.

СМОЛИН В.С.

Смолин В.С.: Как обычно, выскажу противоположное мнение.

Все, что говорилось, основано на простой и всем понятной гипотезе, что всё от Бога, у нас есть душа и нельзя природу описывать чисто научным языком.

Если посмотреть на прогресс человечества, то все постепенно смещается в научное описание.

И когда утверждается, что какой-то академик не знает никаких доказательств того, что можно доверять искусственному интеллекту, то это незнание и вера в Бога с точки зрения науки не являются абсолютным доказательством того, что этого нельзя сделать.

Мы собрались, чтобы говорить не про ИИ, а про знания.

Но почему-то скатились на ИИ, про него можно говорить много.

В принципе, ИИ ничем не отличается от любой другой сферы.

Можно привести пример из банковской деятельности, раз мы собрались на этой площадке. Там в договорах есть мелкий шрифт, который сложно прочесть, и реклама есть, которая не всегда достоверна.

Но почему-то банковскую деятельность никто не думает запрещать из-за этих мелких шероховатостей.

А с искусственным интеллектом все как-то иначе.

Я хотел бы все-таки перейти в обсуждении обратно к знаниям.

У калькулятора они основаны на реализации в технике, а у человека, когда он делает вычисления, знания другие.

Но с позиции конструкторов искусственного интеллекта знания отличаются от информации тем, что это те

данные, которые можно использовать для выполнения каких-то действий.

Соответственно, если посмотреть на развитие наших наук – любых, не только точных, но и гуманитарных, – все науки выделяют в своей деятельности некоторые простые знания: второй закон Ньютона, какие-то философские законы.

И наши навыки состоят в том, что мы применяем эти знания для той конкретной ситуации, которая в реальном мире никогда не повторяется.

Сейчас это в значительной степени проявляется в нейронных сетях.

Кроме того, что у человека есть душа, а у ИИ ее никогда не будет, есть еще второе, вполне научное заблуждение, которое здесь не один раз высказывалось.

О том, что ИИ только комбинирует ранее сделанные открытия.

Это неправда.

Можете взять такой пример.

Некоторые считают, что большие языковые модели собирают цитаты, из которых формируются предложения.

То есть гениальные люди дают какие-то новые мысли, а те, кто обзревает их мысли, своих мыслей не имеют и просто собирают цитаты...

Это не так.

Возьмите любую статью, свою или чужую, выберите из нее пять-шесть слов подряд из любого места и забейте в Яндекс или в Гугл.

Если ваша статья присутствует в Интернете, вам по этим пяти словам ее найдут.

То есть каждые пять-шесть слов в любой статье – не важно, человеком написанных или искусственным интеллектом, – они уникальные.

Так же, как мир наш все время разный, так и тексты, которые совсем не обязательно гениальные, все разные.

Ведь и наши мысли не повторяются, все публикации наших семинаров – разные.

Соответственно, предположение о том, что ИИ не может создавать, потому что у него нет души, мягко говоря, тенденциозно.

Но доказать здесь это никому нельзя.

Тосунян Г.А.: Спасибо.

Смолин В.С.: Просто я хочу обратить внимание, что есть и другое мнение по вопросам, которые мы обсуждаем.

Тосунян Г.А.: Спасибо.

Другое мнение тоже всегда важно. И ценность наших обсуждений именно в том, что высказываются разные точки зрения.

Сергей Федорович, пожалуйста, Вам слово.

Уважаемые коллеги, честно говоря, не хотел участвовать во второй дискуссии, но уж так сильно разгорелся этот полемический пожар, что, в общем-то, нужно немного плеснуть холодной воды.

Так мне кажется.

Или бензина, кто как воспримет.

Дело в том, что мы все время почему-то рассуждаем об ИИ как о некоторой отдельной, разумной, самостоятельной сущности, которая ставит и решает задачи, принимает решения, и эти решения окончательные и пересмотру не подлежат.

И все так опасно, что еще немного – и президент Трамп начнет ядерную войну по совету ИИ, которому он бесконечно доверяет.

Но это эмоции, и они не всегда имеют рациональную основу!

Вместе с тем мы постоянно забываем, что ИИ – это лишь высокотехнологичный инструмент, позволяющий нам эффективно реализовать когнитивно-интеллектуальную деятельность.

И больше ничего!

И мы вправе знать о его возможностях и ограничениях.

Что за продукт породил этот инструмент, можно ли ему доверять, включать в деятельность и так далее?

Главная проблема состоит в квалификации человека, работающего с этим инструментом.

Высококвалифицированный пользователь делает, используя ИИ, продукты эффективные, полезные, удобные и безопасные!

А человек, который не умеет пользоваться этим инструментом, безусловно, малоэффективен и потенциально опасен.

То есть не надо рассматривать существующий искусственный интеллект как отдельную сущность в виде искусственного действующего субъекта, осознающего мир.

Это первое.

Второе.

Мне кажется, что мы очень сильно заикливаемся на вопросе, насколько опасна эта технология для человека и общества.

На самом деле опасность той или иной вещи сосредоточена в руках пользователя.

Скальпелем можно убить и вылечить человека. Важно, в чьих руках этот скальпель.

Мне кажется, что мы с вами в процессе дискуссии еще не дошли до обсуждения действительно важных проблем, которые возникают в силу внедрения технологий ИИ.

По моему мнению, это проблемы интеллектуального симбиоза – эффективного объединения интеллектуальных возможностей человека и ИИ при достижении заданного результата и доверия к информации, генерируемой технологиями.

Только при решении данных проблем мы получаем полезные для общества и личности результаты.

Хотелось бы, чтобы мы отметили именно эти аспекты в проблеме когнитивной безопасности в техногенном мире.

Благодарю!

Тосунян Г.А.: Спасибо.

Пожалуйста, Владимир Сергеевич Кржевов.

к. филос. н. **КРЖЕВОВ В.С.** – акад. ТОСУНЯН Г.А.

Кржевов В.С.: Спасибо.

Я хотел бы продолжить те мотивы, которые сейчас прозвучали в выступлении коллеги Сергеева.

Были произнесены ключевые слова о различии между содержанием информации и ее формой.

Что такое содержание информации?

Что определяет отбор и последующее закрепление того содержания, о котором было сказано?

В конечном счете, это ведь способность действовать сообразно предлагаемым информационным моделям, программам.

Ведь и жизнь, вообще любая, и человеческая деятельность – это информационно направленные процессы.

Информационная направленность – это значит преследование действующим человеком какой-то избранной им цели.

В этой связи мне представляется необходимым всякий раз подчеркивать, что «последней целью» нашей деятельности является самоподдержание, самовоспроизводство.

Информацию, необходимую для этого, мы оцениваем прежде всего по способности послужить нам какими-то моделями, программами, которые могут быть реализованы в нашей предметной вещественно-энергетической деятельности, то есть в производительном поддержании нашего существования.

Искусственный интеллект – это, как было сказано, действительно, средство, инструмент обработки и производства информации, более или менее эффективный, более или менее дееспособный именно в этом плане.

И поэтому его не нужно ни недооценивать, ни переоценивать.

Инструмент этот должен быть оценен с точки зрения своей работоспособности.

А то, что он может быть использован неправильно, то, что он может быть использован во зло...

Опять-таки, зависит от того, в какой системе координат.

Но ведь такие оценки, в сущности, можно отнести ко всему, что создано человеком.

Все его средства, все его орудия могут быть использованы либо для решения задачи поддержания жизни, для минимизации рисков разрушения жизни, либо с противоположной целью.

И вот такие решения принимает, конечно, только человек.

То есть управление, принятие решений и отдача каких-то распоряжений во имя реализации этих решений – это функция человека.

Искусственный интеллект здесь может предложить какую-то программу, какую-то модель, какой-то текст. Но решение – это функция человека.

А решение должно в конечном счете подчиняться задаче сохранения жизни.

Учитывая технологическую вооруженность современной человеческой деятельности, современного высокоспециализированного труда, следовало бы понять, что к нашему теперешнему существованию всецело применима следующая мысль.

Либо мы, нуждаясь друг в друге, научимся между собой договариваться, каким-то образом согласовывать свои интересы и поддерживать стабильные отношения – и тогда все вместе выживем.

Если же не сумеем овладеть этим искусством, то все вместе погибнем.

В современном мире с его технологической вооруженностью господ-победителей быть не может.

Это та истина, к которой, я думаю, нужно прежде всего прийти при обсуждении такого рода проблем.

Спасибо.

Тосунян Г.А.: Спасибо.

Владимир Сергеевич Кржевов в чат написал: «Если я не ошибаюсь, современная теоретическая физика именуется еще и статистической. Это вряд ли можно игнорировать».

Почему Вы так принижаете теоретическую физику? Я что-то не слышал такого термина про современную теоретическую физику.

Кржевов В.С.: Квантовая механика оперирует статистическими данными, если я не ошибаюсь.

Я не специалист по физике, но у меня такое впечатление было всегда.

Тосунян Г.А.: Да, но о статистической физике как замене термина «теоретическая физика» я не слышал. Откровенно говоря, меня это зацепило. Но не будем сейчас это обсуждать.

Кржевов В.С.: У нас же нет возможности оперировать большими массивами данных, не прибегая к методам статистической обработки.

Или я ошибаюсь?

Тосунян Г.А.: Нет, Вы не ошибаетесь, но это не повод назвать теоретическую физику статистической. Это разные вещи.

Кржевов В.С.: С этим я спорить не буду.

Тосунян Г.А.: Хорошо.

Пожалуйста, Александр Викторович Халявкин.

ХАЛЯВКИН А.В.

к. б. н.

У меня несколько реплик.

О термине «социальная инженерия», который, как правило, связывается с мошенничеством и манипуляцией.

А это не так.

Озвучив этот термин, Гарегин Ашотович корректно указал на необходимость в этом случае предварять его словом «криминальная».

Хотя, конечно, есть и некриминальная коннотация этого термина.

Теперь приведу цитату из презентации Михаила Олеговича.

«И 200 лет назад, и сейчас информационный шум, дезинформация находятся на службе у нечестных дельцов и являются способом обмана жалких неудачников, неучей, дураков (таких как Репетилов)».

Это из области негативного.

А область позитивного – это недомолвки, умолчания, ложная активность в дипломатии, в военном деле.

В спорте – это финты в футболе, ложные выпады в фехтовании и тому подобное.

В живой природе – это, например, создание ложного впечатления о себе как о доступной добыче, когда птица, изображая из себя подранка, якобы с трудом перепархивает с места на место и уводит хищника от птенцов.

В науке – это попытки, зачастую удачные, направить вектор исследований по ложному пути, в случае если реализация полученных результатов преждевременна без соответствующей подготовки общества.

Пример из области, которой я активно занимаюсь не один десяток лет, – из геронтологии.

По моему глубокому убеждению, основанному на анализе существующих фактов и тенденций, кардинальный прорыв в полном контроле и отмене старения реален.

Но без соответствующей предварительной подготовки общества к этой реальности возможны негативные последствия, которые хорошо бы предусмотреть и, по крайней мере, минимизировать.

Это как раз и является, в частности, прерогативой истинной социальной инженерии, способной помочь преобразовать инфраструктуру и смысл существования современного общества в гармоничное общество Будущего, то есть осуществить фазовый переход из сущего в должное.

Но этого пока не происходит.

Поэтому, несмотря на существенные вливания, особенно частные, движение в этой области идет в направлении, противоположном заявленной цели.

Благодарю за внимание.

Тосунян Г.А.: Спасибо.

Пожалуйста, Лев Московкин.

МОСКОВКИН Л.И.

Московкин Л.И.: Спасибо большое за это обсуждение.

Я, к сожалению, до сих пор так и не понял, что такое когнитивный, – вопрос, который встал, когда Александр Жуков в интересах «Сбербанка» представлял закон об искусственном интеллекте.

И я так понял из того, что говорил второй докладчик, отвечая на вопросы, что искусственный интеллект – этот хайп или панама, как прозвучало, очень точное определение.

Он необходим для создания шума, для размывания государственного управления.

И прорыв произошел в последние годы, когда программу – возможно, это просто PRISM¹² – научили работать с огромным количеством данных.

Естественно, это все не складывается и потом анализируется, это анализируется на входе.

И в результате появилась возможность – это, конечно, не социальная инженерия, это чисто генетическая инженерия – вербально программировать человека.

Отсюда и терроризм, и мошенничество, и законы, и вообще практически все.

Сам завершающий процесс стал очень дешевым, потому что почти любого человека с помощью генетической вербальной инженерии можно направить буквально куда угодно.

Спасибо всем большое.

Тосунян Г.А.: Спасибо.

Коллеги, будем завершать. Нет возражений?

Хочу еще раз отметить, что все доклады вызвали бурный интерес, и высказать слова благодарности.

Игорь Федорович, сначала Вам заключительное слово.

¹² Платформа для генеративного ИИ и предиктивной аналитики на блокчейне Solana.

ЗАКЛЮЧИТЕЛЬНЫЕ СЛОВА

КЕФЕЛИ И.Ф.

д. филос. н., проф.

Спасибо за предоставленную возможность выступить в заключение очень интересной дискуссии.

Действительно, хайп сам по себе и хайп, связанный с искусственным интеллектом, вызвал очень бурную дискуссию.

Она, на мой взгляд, с одной стороны, весьма интересная, злободневная и своевременная. Но вместе с тем она немножко уводит нас от проблемы человека, от проблемы смысла существования.

Хотя во многих выступлениях проскальзывали такие упреки, что нельзя забывать о смысле нашего существования.

И мне пришла в голову такая шальная мысль – сравнить.

Может ли ИИ воспроизвести героизм – феномен человеческого бытия, человеческой деятельности, человеческой жизни? Личностный героизм, героизм воина на поле сражения, трудовой героизм, массовый героизм и так далее.

Феномен героизма касается самых глубинных свойств человеческого и поведения, и осмысления себя и окружающей действительности. И это, конечно, мы никогда, наверное, не будем сравнивать с тем самым искусственным интеллектом, который, действительно, инструмент в руках человека.

Но инструментом можно пользоваться и во благо, и во зло.

Поэтому все-таки проблема когнитивной безопасности, на мой взгляд, касается в первую очередь безопасности самого по себе человеческого существования.

Лев Московкин в чате написал: «Мне неизвестно, чтобы Россия вела информационную войну. Наоборот, постоянно ругают за пассивность».

Относительно нашей пассивности – можно говорить, что она и есть, и ее нет, то есть эта проблема мало у нас озвучена.

Информационную войну и мы ведем.

Войны не может быть с одной стороны: со стороны Запада или со стороны России.

Или иногда еще пишут про информационную войну против американцев, а не против Америки.

В прошлом году в США вышла книга, в которой утверждалось, что Китай, Россия и Иран ведут информационную войну против американцев. Не против Америки как могущественного государства, а против каждого из американцев.

И это факт реальный.

У меня в группе было несколько китайцев. Я предложил им на практических занятиях перевести эту брошюру и отправить ее в Китай.

Так они сильно удивились. Говорят: неужели такое может быть?

Завершая, скажу спасибо, во-первых, Гарегину Ашотовичу за умелое ведение такого важного мероприятия.

Действительно, многих участников привлекли своевременно поставленные проблемы и когнитивной безопасности, и хайпа.

Мы живем в реальном мире, в котором наша творческая активность должна реализоваться в полной мере во благо нашей страны.

Спасибо.

Тосунян Г.А.: Спасибо.

Пожалуйста, Михаил Олегович.

КОЛБАНЁВ М.О.

д. т. н., проф.

Колбанёв М.О.: Спасибо, Гарегин Ашотович.

Хочу сказать большое спасибо всем: и организаторам этого мероприятия, и всем, кто принял в нем участие.

Самый главный вывод, который можно было бы сделать из нашего заседания.

То, что я услышал, и то, что здесь произошло, – это нисколько не хайп.

О тех проблемах, которые мы сегодня обсуждаем, я услышал высказывания очень содержательные, правильные. Спорные, но научные и обоснованные.

Вспомнилась поговорка: «Поспешишь – людей насмешишь».

Когда у меня стало заканчиваться отведенное для доклада время, я стал торопиться, поспешил и в результате получил критику за невнятно объясненные те или иные свои мысли.

Но даже это получилось хорошо, потому что в подтверждение своей точки зрения я услышал от коллег новые аргументы, которые теперь буду использовать в своей работе.

Еще раз хочу сказать спасибо.

Дай Бог всем здоровья.

Тосунян Г.А.: Спасибо.

Коллеги, у нас практика такая, что даже если докладчик не укладывается в регламент и что-то осталось недовысказанным, то, в принципе, потом в полемике можно это компенсировать.

Михаил Олегович, надеюсь, Вы это тоже почувствовали и в какой-то степени этим воспользовались.

Колбанёв М.О.: Да, совершенно верно.

Тосунян Г.А.: Докладчики всегда потом будут приглашаться на наши заседания в Zoom.

Абдусалам Абдулкеримович, можно Вас послушать?

ЗАКЛЮЧЕНИЕ

ГУСЕЙНОВ А.А.

акад. РАН, д. филос. н., научный руководитель
Института философии РАН

Гусейнов А.А.: Я тоже благодарю докладчиков и всех участников дискуссии.

К сожалению, я заранее не ознакомился с презентациями, но хочу зафиксировать.

Такая практика у нас не первый раз, но сегодня это обозначилось как-то особо.

Это очень хороший способ подготовки наших обсуждений, потому что это самой дискуссии придает характер, приближающийся именно к какому-то академическому разговору.

И по возможности хорошо бы нам и в будущем практиковать это.

Тосунян Г.А.: Это уже давно, это стало системой.

Гусейнов А.А.: Да, это не впервые, но это делается не всегда, я именно это имел в виду.

Обратить на это внимание меня подтолкнуло развернутое выступление профессора Коняевского, который детально прокомментировал доклад Михаила Олеговича и этим придал своему выступлению, я бы сказал, непосредственно академический характер.

Позвольте сделать два замечания.

Разговор вокруг искусственного интеллекта сосредоточился на вопросе безопасности.

Как оградить нас, людей, наши когнитивные способности от тех изменений в окружающей информационной

среде, которые произошли с появлением огромных возможностей искусственного интеллекта?

И как предохранить себя от злонамеренного использования этих возможностей?

Я обратил внимание, что когда у нас идут какие-то качественные сдвиги, появляются качественно новые вещи, связанные с технологиями, то первая общественная реакция, которая выступает на первый план, – запретить, оградить нас от этого.

Когда появился биткоин, какая была реакция? «Да это же вообще какой-то ужас!»

И шла линия на то, чтобы запретить.

Когда появились беспилотные летательные аппараты, когда они впервые начали использоваться, даже в игрушечных форматах, первая реакция была – запретить.

То же самое происходит теперь, когда речь идет об искусственном интеллекте, о таком неконтролируемом «взрыве» и многообразии всего информационного поля, в котором мы живем. Тоже заговорили о том, как нам обеспечить безопасность.

Должен сказать, что это, в общем, правильная постановка вопроса.

Мы живем в мире разнообразных интересов, разнообразных сил – политических, социальных и, как теперь выясняется, цивилизационных.

В этом смысле ограждать свои ценности – это очень важная и, может быть, самая первая и необходимая функция.

Я с этим не спору.

Но как здесь найти такую меру, чтобы это одновременно и не блокировало возможности продуктивного использования потенциала, который несут эти технологии?

Чтобы не получилось так, что мы опять будем потом догонять?

Как это происходило с теми новшествами, о которых я уже говорил: с биткоином, с летательными аппаратами, с тем же ИИ.

Мне кажется, здесь есть над чем подумать.

Обсуждение этого аспекта – важного, хорошего, необходимого аспекта безопасности когнитивных способностей – показало, что на самом деле это такой необходимый шаг, который должен открыть пространство для продуктивного использования возможностей, возникающих здесь.

Своей постановкой вопроса Михаил Олегович как-то нас продвинул, четко обозначив, где границы этой безопасности.

Границы этой безопасности, как он сказал, там, где ИИ выступает как хайп, где он используется в злонамеренных целях.

А это происходит там и тогда, когда не проходит эта граница между человеческим разумом, человеческими когнитивными способностями и искусственным интеллектом.

И совершенно замечательное, простое сравнение, которое нам и не приходит в голову, – про табуретку.

Конечно, она несет на себе следы человека, который создал ее: его способности, умения, его замысел и все что хотите.

Но тем не менее нам не приходит в голову отождествить табуретку и человека.

Между тем в этом отношении нет принципиальной разницы между искусственным интеллектом и табуреткой.

А вот за этими пределами, пожалуйста, уже работайте дальше.

Мы знаем, что табуретку иногда используют во зло: ну, скажем, прибивают человека табуреткой.

Значит ли это, что надо направлять свою энергию на то, чтобы создавать такие табуретки, чтобы их нельзя было использовать в этих побочных деструктивных целях?

И здесь тоже возникает вопрос.

Действительно ли эта граница является такой же безусловной? Хотя и осторожно выражается докладчик: действительно ли она ограничивается разумом человека, его интеллектуальной деятельностью?

Ведь помимо тех строгих расчетов, которые потом оформляются цифровым образом и дают нам искусственный интеллект, в человеческом разуме есть еще что-то другое, невыразимое, то, что потом Гарегин Ашотович назвал душой.

А она что, закрыта для помощи со стороны искусственного интеллекта?

Здесь, действительно, открывается большой простор для фантазии.

Интеллект, разум присущи человеку как живому существу. Они являются способностями органов определенного живого тела.

А почему разум и интеллектуальная сила не могут быть созданы на неорганической основе?

Какие здесь принципиальные запреты?

Дальше.

Искусственный интеллект идет в том направлении, решает те задачи, которые позволяет ему решать человек, которые ставит перед ним человек.

Человек дал ему задание – ИИ в этих рамках и действует.

А почему вы можете исключить, что человек доверит ему и всю полноту своих способностей?

В том числе и творческих, и тех, которые нам, людям, присущи, и сами мы пока не можем даже их внятно выразить.

Если человек увидит, что ИИ сделает это благодаря своим невероятным возможностям намного лучше, чем он сам, чем тогда, когда эти функции базируются на биологической основе нашего брэнного существования...

Если поставить вопрос в таком свободном режиме, то это выглядит не так уж бессмысленно.

Такая свободная постановка вопроса, хотя здесь нельзя заходить слишком далеко, является на самом деле необходимой для того, чтобы повернуться к ИИ с более высоким доверием и не рассматривать каждый случай его использования, интегрирования в нашу жизнь с подозрением, оговорками и так далее.

Конечно, наши аппараты, которые моют окна, полы и делают дома еще какие-то вещи, нельзя путать с собакой и кошкой, с которой мы гуляем, играем и так далее.

Но почему эти аппараты не нуждаются в каком-то особом уходе, уважении и так далее...

То есть это же могут быть другие формы отношения, более живые, теплые и эстетические.

Не оспаривая сам этот рубеж, саму эту границу, что искусственный интеллект – это вторично, это помощник, это средство, это только после человека, нельзя ли допустить какую-то другую эволюционную перспективу?

Давайте вспомним хотя бы идею ноосферы.

Как бы то ни было, мне кажется, сегодня был интересный разговор, так что всем большое спасибо.

Спасибо, что в летние жаркие дни мы все-таки собрались.

Мне было трудно смириться с тем, что каждое заседание начинается ровно в девять утра; точно так же мне трудно смириться, что они не прекращаются в летние месяцы.

Знаю, что президиум РАН прекратил свои заседания до сентября.

Точно так же и университеты, и академические институты.

А наше собрание живет в своем ритме.

Сначала я связывал это с тем, что у Гарегина Ашотовича нет академического опыта и для него летние месяцы такие же рабочие, как и все другие.

Но теперь я вижу, что в этом есть какая-то другая, может быть, более глубокая логика.

И подобно тому, как я все же смирился с тем, что он начинает ровно в девять часов, точно так же я вижу, что и в этой его установке – работать летом – тоже есть какая-то своя логика, что-то правильное.

Но я надеюсь, что все-таки, может быть, заседания будут проходить не обязательно через неделю, а через две-три...

Всего хорошего!

Спасибо, Абдусалам Абдулкеримович, за такую щадящую критику и поддержку.

Логика очень простая. Дело в том, что Академия наук, вузы работают по определенному регламенту, и там ты не можешь не приходить на заседания.

У нас в этом смысле огромное преимущество.

Критерием того, нужно или не нужно проводить заседание, является то, хотят люди приходить, обсуждать, докладывать или не хотят.

И это самое важное основание.

Надеюсь, все наши участники верят, что мы с Абдусаламом Абдулкеримовичем никого насильно сюда не приглашаем, не принуждаем к участию или выступлениям.

Да, летом, конечно, может быть немного иной регламент. Мы через две недели собираемся встретиться, но в августе, скорее всего, будет уже трехнедельный перерыв.

Будет очно-заочное заседание. Впрочем, как жизнь сложится... Для нас важно максимально давать возможность людям поучаствовать в обсуждении вопросов, которые нас интересуют.

Выскажу несколько коротких тезисов по обсуждаемым сегодня вопросам, и на этом завершим.

Конечно, начну еще раз с благодарности нашим докладчикам и содокладчикам, потому что Валерий Аркадьевич тоже привнес свою лепту в обсуждение, обостренно высказав свои возражения.

Игорь Федорович и Михаил Олегович сделали основные доклады. И все с интересом их слушали, потому что в них были затронуты фундаментальные вопросы познания.

Тем более что когнитология – это междисциплинарная область знания, и на нашем междисциплинарном заседании это приобретает особое звучание.

Абдусалам Абдулкеримович сказал, что наличие заранее присланных презентаций придает научность обсуждению.

Мне кажется, что в подавляющем большинстве случаев, даже когда заранее докладчики не успевают или не считают нужным присылать свои материалы, тем не менее упрекнуть нас в обсуждении ненаучном вряд ли будет справедливо. Хотя междисциплинарность нашего Совета вносит свои коррективы в характер обсуждения.

Мы стараемся анализировать проблемы, не уходя слишком вглубь отрасли, а делая акцент на более широком их осмыслении. Потому что математики вряд ли могут углубиться в тонкости философских вопросов, а философы – в тонкости биологических или медицинских проблем, которые мы тоже обсуждаем.

На междисциплинарном уровне, чтобы состоялось обсуждение, важно, чтобы была обеспечена доступность, понятность темы не только для узких специалистов (для них есть специальные семинары, научные советы).

А преимущество нашего Совета именно в том, что, может быть, не слишком глубоко, иногда обобщенно рассматриваются отдельные темы. Но в этом нельзя видеть ненаучность. В этом и есть особенность междисциплинарного обсуждения.

В данном случае я пытаюсь полемизировать сам с собой. Объясняя, почему мы все-таки на таком доступном, простом, понятном языке стараемся выносить на обсуждение глубокие темы.

Огромная благодарность докладчикам за то, что, когда они слышат вопросы людей из другой отрасли, они не упрекают их в дилетантизме, а пытаются понять вопрос и ответить. В этом наше большое преимущество.

Был поднят вопрос о том, что информация может или не может быть ложной, и о том, что дезинформация – это передача ложной информации.

Кто-то из докладчиков уточнил: так все-таки есть дезинформация?

А как иначе дезинформацию расшифровать, кроме как передачу ложной информации?

На самом деле понятия «знание», «информация» – это базовые понятия.

И мы предложили Галине Вениаминовне Сориной сделать доклад на нашем «Рабочем завтраке» об этих базовых понятиях.

Хотел бы также обратить внимание на мысль, высказанную на заседании, о том, что 10-15% дезинформации могут погубить систему.

Обычно эффективно функционирует система, основанная на доверии.

А доверие, конечно, предполагает минимум лжи, обмана и дезинформации (я не знаю, как этот термин использовать аккуратнее).

И если мы будем обсуждать тему об информации, знании, то надо будет вернуться к вопросу о том, насколько допустим, насколько возможен уровень погрешности, при которой система будет работать эффективно, а не разрушаться.

Маленькая реплика по поводу того, всегда ли 5 плюс 6 равно 11. Или сумма может превратиться в яичницу, если это 5 помидоров и 6 яиц?

Это один из тех вопросов, которые тоже требуют осмысления, хотя бы на таком, может быть, бытовом, банальном уровне.

Но 5 плюс 6, если это элементы системы, – не важно, это яйца, масло, сковородки, – все равно это 11 элементов системы.

Если в узком смысле рассматривать их как яйца и помидоры, тогда можно говорить о яичнице. Кстати, к яичнице надо еще приложить определенные действия.

Поэтому 5 и 6 – это не просто 11. Это всегда 11 элементов системы, но важно понять, что в этой системе происходит.

И наконец, прозвучала короткая реплика по поводу сопоставления человека с искусственным интеллектом.

Была высказана мысль, что это уже хайп.

Я придерживаюсь той же позиции, что и Абдусалам Абдулкеримович. Он задал вопрос: почему разум не может быть создан на неорганической основе и реализован на органической основе?

Наверное, меня можно упрекнуть в том, что я порой ухожу от сугубо научного подхода и ссылаюсь на Бога, на душу, на эти категории.

Но если исходить из того, что материя первична, а сознание, душа и все остальное – вторично, тогда да, в материальном мире интеллект – это производная от материи, которая называется человеком.

Дело в том, что я уверен: материя отнюдь не первична, и тем более душа – не вторична. Это сочетание дает тот многообразный мир, который мы познаем и никогда до конца не познаем, пока находимся в материальной сфере бытия.

В какой ипостаси мы сегодня и здесь участвуем, и проживаем эту жизнь?

Эта материальная часть отнюдь не является сутью мироздания и тем более вершиной бытия.

Божественное начало делает все-таки человека, человеческое сознание, человеческую душу чем-то более высоким.

Конечно, в научном обсуждении это может звучать, немного резонируя с мнением людей, приверженных исключительно научным стандартам мышления. Но у всех, что называется, есть свои недостатки.

Если быть абсолютным материалистом, то, наверное, можно предположить, что сравнение человека и искусственного интеллекта – это не хайп. Это вполне несет в себе какие-то риски, возможности в определенном смысле замены и даже подчинения человека.

Но хотел бы поделиться со всеми моей уверенностью в том, что абсолютно невозможно, чтобы искусственный интеллект заменил человека вне тех пределов, в которых человек этого захочет.

Так что играть в шахматы – это одно. А способность человеческого мышления и интеллектуальные навыки, таланты, которые присущи исключительно мыслящему и чувствующему человеку, – это другое.

Наука не все может объяснить. Надо признать, что она ограничена в своих возможностях.

И не является ненаучностью признание того, что мы не всё можем объяснить исключительно теми научными знаниями, которыми мы сегодня владеем.

На этом я остановлюсь, а то слишком много критики посыплется в мой адрес со стороны моих уважаемых коллег.

Если никто замечаний не имеет после моих рассуждений, на этом закончим наше заседание.

Спасибо.

**Список литературы, опубликованной по итогам
заседаний НКС ООН РАН, НИИ ДДиП
и «Открытых дискуссий» президента АРБ**

1. Анализируя сегодня, говорим и думаем о будущем (18.04.2020) / под общ. ред. академика РАН Г.А. Тосуняна. – М.: ООО «Новые печатные технологии», 2022. – 175 с.
2. Предновогодние заседания: 2020, 2021 (31.12.2020, 31.12.2021) / под общ. ред. академика РАН Г.А. Тосуняна. – М.: ООО «Новые печатные технологии», 2022. – 206 с.
3. Ответственность пациентов и врачей. Уровень здравоохранения в России (03.04.2021) / под общ. ред. академика РАН Г.А. Тосуняна. – М.: ООО «Новые печатные технологии», 2022. – 124 с.
4. Конкурентоспособность российской науки: проблемы и решения (03.04.2021, 17.04.2021, 15.05.2021) / под общ. ред. академика РАН Г.А. Тосуняна. – М.: ООО «Новые печатные технологии», 2022. – 333 с.
5. О проекте «Стратегия развития финансового рынка до 2030 года» (09.10.2021) / под общ. ред. академика РАН Г.А. Тосуняна. – М.: ООО «Новые печатные технологии», 2021. – 155 с.
6. О развитии конкуренции в сфере науки (30.10.2021) / под общ. ред. академика РАН Г.А. Тосуняна. – М.: ООО «Новые печатные технологии», 2022. – 130 с.
7. Социально-профессиональные проблемы прекаризации труда (18.12.2021) / под общ. ред. академика РАН Г.А. Тосуняна. – М.: ООО «Новые печатные технологии», 2023. – 131 с.

8. Инвалидность и жизнь (12.02.2022) / под общ. ред. академика РАН Г.А. Тосуняна. – М.: ООО «Новые печатные технологии», 2023. – 106 с.
9. Новая экономическая реальность: региональный разрез. Российский рынок драгоценных металлов (21.04.2022, 15.10.2022) / под общ. ред. академика РАН Г.А. Тосуняна. – М.: ООО «Новые печатные технологии», 2023. – 161 с.
10. 1. Санкции. 2. Перспективы экспорта российских нефти и газа в условиях санкционного давления. 3. Интернет-торговля: текущая ситуация и перспективы (11.06.2022, 25.06.2022) / под общ. ред. академика РАН Г.А. Тосуняна. – М.: ООО «Новые печатные технологии», 2022. – 242 с.
11. Демография России: тренды последних лет и краткосрочный прогноз (15.10.2022) / под общ. ред. академика РАН Г.А. Тосуняна. – М.: ООО «Новые печатные технологии», 2024. – 120 с.
12. Общее образование: проблемы и решения (29.10.2022) / под общ. ред. академика РАН Г.А. Тосуняна. – М.: ООО «Новые печатные технологии», 2023. – 148 с.
13. Китай: вчера, сегодня, завтра (19.11.2022) / под общ. ред. академика РАН Г.А. Тосуняна. – М.: ООО «Новые печатные технологии», 2023. – 189 с.
14. Одаренные дети. «Гадкие лебеди» братьев Стругацких как антиутопия кризиса образования: межпоколенческий дефолт (17.12.2022) / под общ. ред. академика РАН Г.А. Тосуняна. – М.: ООО «Новые печатные технологии», 2023. – 163 с.
15. 2022-й – «особый». 2023-й – риски и ожидания (31.12.2022) / под общ. ред. академика РАН Г.А. Тосуняна. – М.: ООО «Новые печатные технологии», 2024. – 134 с.

16. Закат общества конкуренции и коллаборативное преимущество (21.01.2023) / под общ. ред. академика РАН Г.А. Тосуняна. – М.: ООО «Новые печатные технологии», 2023. – 128 с.
17. 1. Мировой океан: ресурсы и влияние на климат. 2. Безусловный базовый доход: шанс для России? (04.02.2023) / под общ. ред. академика РАН Г.А. Тосуняна. – М.: ООО «Новые печатные технологии», 2023. – 148 с.
18. Психологическое состояние российского общества (18.03.2023) / под общ. ред. академика РАН Г.А. Тосуняна. – М.: ООО «Новые печатные технологии», 2023. – 192 с.
19. О мозге (01.04.2023) / под общ. ред. академика РАН Г.А. Тосуняна. – М.: ООО «Новые печатные технологии», 2023. – 187 с.
20. Китай: открытая дискуссия. Социальный рейтинг в Китае (26.04.2023, 27.05.2023) / под общ. ред. академика РАН Г.А. Тосуняна. – М.: ООО «Новые печатные технологии», 2024. – 185 с.
21. Индия: вчера, сегодня, завтра. Взаимодействие России и Индии в условиях глубокой структурной трансформации российской экономики (29.04.2023) / под общ. ред. академика РАН Г.А. Тосуняна. – М.: ООО «Новые печатные технологии», 2023. – 152 с.
22. Денежно-кредитная политика и монетизация экономики (13.05.2023, 11.05.2024) / под общ. ред. академика РАН Г.А. Тосуняна. – М.: ООО «Новые печатные технологии», 2024. – 238 с.
23. Молодежь и мошенничество (31.05.2023). Теория поколений и модели мира (22.06.2023) / под общ. ред. академика РАН Г.А. Тосуняна. – М.: ООО «Новые печатные технологии», 2023. – 158 с.

24. Социальное неравенство (10.06.2023) / под общ. ред. академика РАН Г.А. Тосуняна. – М.: ООО «Новые печатные технологии», 2023. – 145 с.
25. Национальная сила: оценка и практическое применение. Гипотеза общественного прогресса: аргументы «за» и «против» (24.06.2023) / под общ. ред. академика РАН Г.А. Тосуняна. – М.: ООО «Новые печатные технологии», 2024. – 179 с.
26. Турция: вчера, сегодня, завтра (08.07.2023) / под общ. ред. академика РАН Г.А. Тосуняна. – М.: ООО «Новые печатные технологии», 2025. – 102 с.
27. Научное лидерство и человеческий капитал (22.07.2023) / под общ. ред. академика РАН Г.А. Тосуняна. – М.: ООО «Новые печатные технологии», 2024. – 150 с.
28. Цифровые валюты центральных банков (26.08.2023) / под общ. ред. академика РАН Г.А. Тосуняна. – М.: ООО «Новые печатные технологии», 2024. – 151 с.
29. Общество и государство (09.09.2023) / под общ. ред. академика РАН Г.А. Тосуняна. – М.: ООО «Новые печатные технологии», 2024. – 163 с.
30. Искусственный интеллект (14.10.2023) / под общ. ред. академика РАН Г.А. Тосуняна. – М.: ООО «Новые печатные технологии», 2024. – 182 с.
31. «Зеленая» экономика: принципы и проблемы (18.11.2023) / под общ. ред. академика РАН Г.А. Тосуняна. – М.: ООО «Новые печатные технологии», 2024. – 164 с.
32. Институт финансового омбудсмена, его роль в развитии общества (25.11.2023) / под общ. ред. академика РАН Г.А. Тосуняна. – М.: ООО «Новые печатные технологии», 2024. – 152 с.

33. 2023-й – итоги. 2024-й – перспективы (30.12.2023) / под общ. ред. академика РАН Г.А. Тосуняна. – М.: ООО «Новые печатные технологии», 2025. – 136 с.
34. В.И. Вернадский и его наследие (20.01.2024) / под общ. ред. академика РАН Г.А. Тосуняна. – М.: ООО «Новые печатные технологии», 2026. – 175 с.
35. Генномодифицированные продукты: «за» и «против» (03.02.2024) / под общ. ред. академика РАН Г.А. Тосуняна. – М.: ООО «Новые печатные технологии», 2025. – 102 с.
36. Банки Китая: стратегия развития (03.02.24) / под общ. ред. академика РАН Г.А. Тосуняна. – М.: ООО «Новые печатные технологии», 2024. – 90 с.
37. Проблема общечеловеческих ценностей. Причины ценностных противостояний в современном мире (17.02.2024) / под общ. ред. академика РАН Г.А. Тосуняна. – М.: ООО «Новые печатные технологии», 2024. – 162 с.
38. Искусственный интеллект в банковской сфере (29.02.2024) / под общ. ред. академика РАН Г.А. Тосуняна. – М.: ООО «Новые печатные технологии», 2024. – 121 с.
39. Малый и средний бизнес (02.03.2025) / под общ. ред. академика РАН Г.А. Тосуняна. – М.: ООО «Новые печатные технологии», 2025. – 172 с.
40. Мозговая активность в пожилом возрасте (паркинсон, альцгеймер, деменция) (06.04.2024) / под общ. ред. академика РАН Г.А. Тосуняна. – М.: ООО «Новые печатные технологии», 2024. – 162 с.
41. БРИКС: платежные системы (20.04.2024) / под общ. ред. академика РАН Г.А. Тосуняна. – М.: ООО «Новые печатные технологии», 2024. – 173 с.

42. Цивилизация. Что это? (25.05.2024) / под общ. ред. академика РАН Г.А. Тосуняна. – М.: ООО «Новые печатные технологии», 2025. – 208 с.
43. Жилая недвижимость (27.07.2024) / под общ. ред. академика РАН Г.А. Тосуняна. – М.: ООО «Новые печатные технологии», 2025. – 166 с.
44. Общественный договор (07.09.2024) / под общ. ред. академика РАН Г.А. Тосуняна. – М.: ООО «Новые печатные технологии», 2025. – 186 с.
45. Банки и конкуренция: материалы «Открытой дискуссии президента АРБ» (19.09.2024) / под общ. ред. академика РАН Г.А. Тосуняна. – М.: ООО «Новые печатные технологии», 2024. – 100 с.
46. Жизнь и смерть. Право на эвтаназию и суицид (21.09.2024) / под общ. ред. академика РАН Г.А. Тосуняна. – М.: ООО «Новые печатные технологии», 2025. – 221 с.
47. Наука и власть (05.10.2024) / под общ. ред. академика РАН Г.А. Тосуняна. – М.: ООО «Новые печатные технологии», 2025. – 206 с.
48. Борьба с кризисом и промышленная политика (19.10.2024) / под общ. ред. академика РАН Г.А. Тосуняна. – М.: ООО «Новые печатные технологии», 2025. – 178 с.
49. Философия в современном мире (16.11.2024) / под общ. ред. академика РАН Г.А. Тосуняна. – М.: ООО «Новые печатные технологии», 2025. – 218 с.
50. Конкуренция на финансовом рынке (30.11.2024) / под общ. ред. академика РАН Г.А. Тосуняна. – М.: ООО «Новые печатные технологии», 2025. – 133 с.

51. Роль институтов в развитии общества (14.12.2024) / под общ. ред. академика РАН Г.А. Тосуняна. – М.: ООО «Новые печатные технологии», 2025. – 211 с.
52. Инфляция у каждого своя: материалы «Открытой дискуссии президента АРБ» (17.12.2024) / под общ. ред. академика РАН Г.А. Тосуняна. – М.: ООО «Новые печатные технологии», 2025. – 108 с.
53. Завершился 2024 год. Что дальше? (31.12.2024) / под общ. ред. академика РАН Г.А. Тосуняна. – М.: ООО «Новые печатные технологии», 2025. – 97 с.
54. Что такое Свобода? Свобода воли, свободы личности (15.02.2025) / под общ. ред. академика РАН Г.А. Тосуняна. – М.: ООО «Новые печатные технологии», 2025. – 210 с.
55. Цивилизационные основы России (часть 2): Качество жизни. Равноправие женщин и мужчин. Русское крестьянство (01.03.2025) / под общ. ред. академика РАН Г.А. Тосуняна. – М.: ООО «Новые печатные технологии», 2026. – 178 с.
56. Коррупция в Китае (22.03.2025) / под общ. ред. академика РАН Г.А. Тосуняна. – М.: ООО «Новые печатные технологии», 2026. – 143 с.
57. Религия и общество (05.04.2025) / под общ. ред. академика РАН Г.А. Тосуняна. – М.: ООО «Новые печатные технологии», 2026. – 164 с.
58. Многополярный мир: Россия, США, Китай, Индия, ЕС: влияние на него новой администрации США (31.05.2025) / под общ. ред. академика РАН Г.А. Тосуняна. – М.: ООО «Новые печатные технологии», 2026. – 167 с.

59. Редактирование генома человека (25.10.2025) / под общ. ред. академика РАН Г.А. Тосуняна. – М.: ООО «Новые печатные технологии», 2026. – 178 с.

Электронные версии сборников
можно скачать по ссылке:

<https://rannks.ru>

Когнитивная безопасность

Материалы заседания
НКС ООН РАН и НИИ ДДиП
12 июля 2025 года

Презентации докладчиков
можно скачать по ссылке:

<https://rannks.ru/pubs/>

Подписано в печать 05.08.2026
Формат 60х90/16
Цифровая печать
Тираж 500 экз. Заказ № 84

Отпечатано в ООО «НОВЫЕ ПЕЧАТНЫЕ ТЕХНОЛОГИИ»
117525, г. Москва, ул. Днепропетровская, д. 3, корп. 5, пом. III

Научно-консультативный совет Отделения общественных наук РАН был создан в 2012 году как Совет по правовым, экономическим, социально-политическим и психологическим аспектам финансово-кредитной системы. В феврале 2020 года члены НКС приняли решение расширить компетенцию Совета, перейдя от рассмотрения вопросов развития финансового рынка к более широкому кругу проблем развития общества, поставив во главу угла своих исследований и дискуссий вопросы: в каком обществе мы живем? Какое общество мы хотели бы оставить своим потомкам в наследство?

Сопредседатели Совета: академики РАН А.А. Гусейнов, А.А. Кокошин и Г.А. Тосунян.

Ассоциация российских банков учреждена в марте 1991 года. Миссия Ассоциации российских банков – реализация программы банкизации страны, создание условий для эффективного функционирования, развития банковской системы России и обеспечения ее стабильности, защиты прав, интересов банков и условий для справедливой рыночной конкуренции; участие в построении национальной финансовой экосистемы, основанной на принципах соблюдения прав и реализации комплекса мер по повышению финансовой грамотности потребителей.

Национальный исследовательский институт Доверия, Достоинства и Права учрежден в конце 2019 года.

Цель института – многогранное изучение вопросов человеческой жизнедеятельности и общественных процессов, которые наибольшим образом влияют на развитие доверия в обществе, повышение чувства собственного достоинства у граждан страны и на формирование уважения друг к другу.

Институт приступил к работе в начале 2020 года в формате научных заседаний с коллегами, интересующимися проблемами доверия, достоинства, их правового обеспечения и стимулирования.