

АНДРЕЙ
КУРПАТОВ

АВТОР БЕСТСЕЛЛЕРА
«КРАСНАЯ ТАБЛЕТКА»



БУДУЩЕЕ
УЖЕ РЯДОМ

ЧЕТВЁРТАЯ МИРОВАЯ ВОЙНА

КНИГА ДЛЯ ИНТЕЛЛЕКТУАЛЬНОГО МЕНЬШИНСТВА

АКАДЕМИЯ
СМЫСЛА

Академия смысла

Андрей Курпатов

**Четвертая мировая
война. Будущее уже рядом**

«Курпатов А.В.»

2018

УДК 159.922.1
ББК 88.53

Курпатов А. В.

Четвертая мировая война. Будущее уже рядом / А. В. Курпатов —
«Курпатов А.В.», 2018 — (Академия смысла)

ISBN 978-5-6040992-5-4

В ближайшие десятилетия мир переживёт самую значительную трансформацию за всю историю человечества. Технологии радикально изменят политику и экономику, среду обитания и отношения между людьми. Изменимся и мы сами. До неузнаваемости. Эта книга расскажет о том, почему искусственный интеллект — не выдумка, о том, как это работает, и почему он лучше наших мозгов. Вы узнаете, как он думает, и к каким последствиям приведут новейшие научные открытия. Вас ждут все сценарии возможного будущего... Поможет ли это вам подготовиться к новой реальности? Нет. Но вы серьезно задумаетесь о том, что происходит сейчас!

УДК 159.922.1

ББК 88.53

ISBN 978-5-6040992-5-4

© Курпатов А. В., 2018

© Курпатов А.В., 2018

Содержание

Глава первая	9
Да будет так!	10
Повсеместная роботизация	14
Взбесившийся 3D-принтер	21
Навстречу сингулярности	32
Время не ждёт	37
Глава вторая	41
Не в нашу пользу	42
Языковая игра	48
Конец ознакомительного фрагмента.	51

Доктор Андрей Курпатов

Четвёртая мировая война.

Будущее уже рядом!

С надеждой автор посвящает эту книгу всем студентам «Академии смысла».

Не нужно быть сверхразумным искусственным интеллектом, чтобы понять: двигаться навстречу величайшему событию в истории человечества и не готовиться к этому – просто глупо.

МАКС ТЕГМАРК, Массачусетский технологический институт

Книга для интеллектуального меньшинства

абсолютно не рекомендована тем, кто готов по любому поводу оскорбиться

Наша главная слабость – в глупости, лени и самодовольстве.

В 2016 году я начал публиковать на портале «Сноб» цикл статей под общим названием «Четвёртая мировая». Они были посвящены нашему скорому и отнюдь не безоблачному будущему.

На рубеже веков человечество оказалось перед лицом новой реальности: «третья информационная волна» (Элвин Тоффлер), «четвёртая технологическая революция» (Клаус Шваб), «технологическая сингулярность» (Рэй Курцвейл).

То есть наша цивилизация трансформируется, причём фундаментальным образом. Но что мы знаем о рисках, о возможных последствиях этих перемен? Задумываемся ли мы о них всерьёз?

Мои статьи приняли тогда, в целом, положительно – сотни тысяч просмотров, много добродетельных отзывов. Однако был и весьма характерный фон, я бы даже сказал – «душок». Мол, всех тут «доктор из телевизора» запугивает, а никакой угрозы нет: технологии, информационный бум и искусственный интеллект – это всё прекрасно, и нечего паниковать.

Кто-то говорил, что мои «пророчества» – дело столь отдалённого будущего, что даже нелепо об этом думать. Кто-то утверждал, что в реальном мире программистов и специалистов по искусственному интеллекту «всё вообще по-другому» и нечего тут «психологам» лезть с футуристическими прогнозами. Кто-то утверждал, что я и вовсе ретроград, луддит, противник прогресса и цивилизации.

Но считать меня луддитом столь же нелепо, как и называть тем самым «психологом» (я всё-таки врач-психиатр, что далеко не одно и то же). Новые технологии – это замечательно, я и правда так считаю. Однако предельно глупо, на мой взгляд, развивать технологии, которые в корне меняют среду нашего обитания, не учитывая возможные последствия для продукта этой самой среды – то есть для нас с вами.

Мы – плоть от плоти – та среда, которая нас окружает: и не только физико-химическая, но и языковая, культурная, психологическая, идеологическая – то есть собственно информационная.

В нас нет ничего «своего», мы полностью сделаны из окружающей нас среды. Допускаю, что это тяжело принять тем, кто верит в «духовный рост», «божественный замысел» и проповедует «любовь к себе», но такова правда.

- На физическом уровне мы то, что мы физически потребляем: химические вещества, находящиеся в пище, воде, вдыхаемом воздухе (грубо говоря, мы то, что мы едим, что пьём, чем дышим).

- На информационном уровне мы являемся производными той информационной среды, в которой живём, – воспитание и образование, поведенческие стереотипы в обществе, массмедиа.

Но в чём тут, вы скажете, новость?.. Если взглянуть на историю человечества, то информационная среда менялась регулярно, зачастую радикально – и никаких проблем! С чего бы им теперь вдруг возникнуть?

Да, менялась, но раньше эти изменения касались только содержания – трансформировались представления людей о мире, эволюционировали культурные паттерны и т. д. Сейчас же изменяется сама *структура* информационной среды.

Причём подобные структурные «фазовые переходы» человечество уже переживало – изобретение письменности, печатного станка, телеграфа, радио, синемаатографа. И за подобными «переходами» всегда следовала, по сути, новая эра в истории человечества.

Но посмотрите, как эти эпохи ужимаются: от момента появления письменности до печатного станка – тысячи лет, от станка до телеграфа – сотни, дальше – десятки.

Сейчас новые способы распространения информации появляются чуть ли не каждый год: интернет, электронная почта, интернет-поисковики, мобильный интернет, социальные сети и т. д., и т. п.

Можно с полной уверенностью утверждать, что ещё никогда за всю историю человечества структурные изменения в информационном поле не были столь грандиозными и значительными, как сейчас.

Информационные технологии, роботизация и уберизация, а также собственно искусственный интеллект превращаются в своеобразный экзоскелет нашего мозга, а это естественным образом приводит к неизбежной атрофии интеллектуальной функции.

С мозгами как с мышцами: если их функцию выполняет какой-то сторонний агрегат, то они медленно, но верно усыхают.

Из-за социальных сетей, эффекта постоянной подключённости («всегда на связи»), агрессивной конкуренции между производителями контента, цифровой зависимости и других новых «зол» изменилось не только количество, но и качество потребляемой нами информации.

Эта фундаментальная трансформация среды с неизбежностью приводит к нашим собственным изменениям. Но из-за когнитивных искажений мы субъективно занижаем значение происходящего: к переменам мы стали привыкать быстро, а собственных изменений не видим, потому что не с чем сравнить – всё человечество меняется разом.

Многие, впрочем, чувствуют: «что-то пошло не так». Изменения вроде бы и положительные, но вот фон – нет, какой-то странный: всё сложнее определиться с целями, жизненные перспективы выглядят какими-то туманными (если вообще просматриваются), нарастает чувство безысходности, отношения между людьми становятся всё более и более поверхностными и формальными.

«Технологии будут систематически менять наше понимание того, что значит быть человеком, что значит быть в социуме и что значит заниматься политикой. [...] Мы действительно проходим через сдвиг парадигмы. Она замечательная всем тем, что нам даёт, но одновременно ведёт и к ненадёжности существующих структур, которые теряют свою

ценность и значение. Следовательно, этот новый режим бытия требует нового мирового порядка».

НИШАН ШАХ, Центр цифровой культуры Люнебургского университета

Отражают ли эти смутные ощущения действительный масштаб перемен? Сомневаюсь. Да и вопросов больше, чем ответов... Мы до сих пор не понимаем, в чём, собственно, эти изменения заключаются, что будет с нами дальше, как изменится наше общество.

В любом случае просчёт возможных рисков, связанных с технологическим и цифровым «улучшением жизни», – это важная задача.

МЕДИЦИНСКИЙ ПРИМЕР

В своё время мы вмешались в естественный отбор, спасая жизни людей при помощи антибиотиков и обезболивающих при хирургических операциях. Мы хорошо лечим рак, активно развиваются протезирование и трансплантология.

Невероятные успехи достигнуты в экстракорпоральном оплодотворении, сохранении беременности и неонатальной медицине, детская смертность стала минимальной.

Современные нейролептики и антидепрессанты позволяют лицам, страдающим психическими расстройствами, вести полноценную жизнь.

Успех просто невероятный: на планете сейчас живёт больше людей, чем за всю её историю, а средняя продолжительность жизни человека лишь за один прошлый век увеличилась более чем в два раза.

Но само это благоденствие вызывает проблемы, которые пока непонятно как решать: супербактерии, рост патогенности вирусов и появление новых¹, рост психических расстройств и врождённых патологий. И это, конечно, далеко не полный список...

Тех из нас, кого эволюция раньше бы выбраковала, современная медицина спасает. В геноме человечества происходит накопление предрасположенностей к самому широкому кругу болезней. И потому уже сейчас рождение ребёнка без патологий и более-менее устойчивого к болезням – что-то за гранью фантастики.

Да, достижения медицины – это замечательно (меня они особенно радуют, ведь я бы уж точно давно оказался в числе выбракованных эволюцией особей). Но есть у этой медали и обратная сторона.

Врачи думают о последствиях своего вмешательства в естественный отбор. Они осознают риски и с удвоенной силой занимаются вопросами вирусологии, иммунологии и генной терапии. Но я не видел никого, кто был бы настолько же всерьёз озабочен последствиями фундаментальной трансформации информационной среды.

Есть единичные исследователи, которые открыто говорят о возможных рисках, но их голоса, к сожалению, или игнорируются, или не выглядят достаточно убедительными. А общая реакция общества и различных его институтов вполне укладывается в формулу, с которой я начал: глупость, леность и самодовольство.

Прошло не так много времени с публикации того моего «снобовского» цикла статей, а количество «критиков» уже существенно поубавилось.

¹ Вот лишь небольшой их список: ВИЧ, заражение людей птичьим гриппом, геморрагические лихорадки (Эбола и др.), новые разновидности вирусного гепатита и т. д.

То, что казалось каким-то совершенно отдалённым будущим: машины-беспилотники, 3D-принтеры, позволяющие работать практически с любыми материалами, чипы в человеческих головах, детальная персонализация человека по его поведению в сети и т. д., – всё это уже, так сказать, в дверях.

Речь идёт не о каких-то «нюансах», а о системной проблеме: **перед нами не только собственно технологические риски, но и экономические, общественно-политические, экзистенциальные.**

- Технологические риски связаны прежде всего с возможностью неконтролируемого развития искусственного интеллекта.

- Экономические риски связаны с массовой безработицей, обусловленной полной автоматизацией производства, что приведёт к системному кризису современной модели экономики.

- Общественно-политические риски – это и возможная кибервойна, и возникновение тоталитарных государств (квазигосударств), управляемых собственниками BigData.

- Экзистенциальные риски в грядущем цифровом мире связаны с утратой человечности в традиционном её понимании, а также с интеллектуальной деградацией общества.

Каждое из этих направлений разрабатывается независимыми экспертами, в университетской среде и исследовательскими компаниями. Идёт активная дискуссия, но общей картины пока, к сожалению, нет.

В этой книге я постараюсь рассказать о проблемах, связанных с наступлением «четвёртой промышленной революции», торжественно провозглашённой на Давосском экономическом форуме его бессменным президентом Клаусом Швабом.

Да, когда представители гигантского транснационального бизнеса самозабвенно рассказывают нам о грядущем счастье, я предпочитаю говорить о реальности. Мы должны оценить, насколько указанные риски взаимосвязаны и какова вероятность, что они вызовут эффект домино.

Ну и, конечно, я добавлю к этому скорбному списку свою «ложку дёгтя». Даже применяя на себя роль футуролога, я не могу перестать быть врачом-психиатром, а на мой профессиональный взгляд, самой серьёзной проблемой нового времени будет деформация психики человека.

Этому аспекту, этому «слабому звену» обычно уделяется совсем мало внимания, но именно это «звено», как мне кажется, и запустит ту самую цепочку падающих друг на друга костяшек домино.

Но обо всех костяшках по порядку...

Глава первая Цифровой рай

Лучший способ предсказать будущее – это изобрести его.
АЛАН КЭЙ

В этой главе мы рассмотрим скорое будущее, которое нам обещает главный пророк современных технологий – Рэй Курцвейл.

Рэй Курцвейл – личность, без преувеличения, легендарная. С победами на поприще информатики его поздравляли президенты США – Линдон Джонсон (Рей было тогда 20 лет от роду) и Билл Клинтон, вручивший Курцвейлу в 1999 году информационного Нобеля – National Medal of Technology.

Курцвейл создал первый музыкальный синтезатор, первый планшетный сканер, первую читающую машину для слепых, первым научил компьютеры распознавать человеческую речь. И это только некоторые из его личных достижений, не считая работы на Google, IBM ит. д.

Сейчас Курцвейл работает техническим директором Google, где возглавляет все работы по искусственному интеллекту. А в качестве хобби создаёт помощника, «способного отвечать на наши вопросы ещё *до того*, как вы их сформулируете». Нет, я не шучу. Это цитата.

Впрочем, Рэй Курцвейл, конечно, более известен широкой общественности как футуролог. В книге «Эпоха духовных машин» он сформулировал «закон ускоряющейся отдачи», который позволяет ему с удивительной точностью предсказывать – буквально по годам – достижения в области развития компьютерных технологий и искусственного интеллекта.

Да будет так!

Я придумал закон ускоренной отдачи, чтобы правильно рассчитывать время в моих собственных технологических проектах: чтобы я мог начинать их за несколько лет до того, как они станут осуществимыми.

РЭЙ КУРЦВЕЙЛ

Согласно закону ускоряющейся отдачи, развитие технологий происходит экспоненциально: то есть чем мощнее становится та или иная технология, тем большее ускорение в своём развитии она приобретает.

«За семь лет проект “Геном человека”² собрал один процент генома, – рассказывает Курцвейл. – Мейнстримовые критики заявляли: “Я же говорил, что ничего не получится. За семь лет – один процент, значит, на весь геном уйдёт 700 лет”. Моя реакция была другой: “Ого, мы уже сделали один процент? Мы почти закончили!” Дело в том, что один процент – это всего семь удвоений до ста процентов. Удвоение происходит каждый год. И действительно, проект закончили уже через семь лет. То же самое произошло со стоимостью: первый геном стоил миллиард долларов, а сейчас эта процедура стоит всего 1000 долларов».

Чтобы представить себе, о чём говорит Курцвейл, вспомните знаменитую притчу о создателе шахматной игры – бедном мудреце и математике Сета.

Шахматы так впечатлили индусского царя Шерама, что он решил беспримерно наградить Сета.

– Я настолько богат, что могу исполнить любое твоё самое смелое желание, – сказал царь мудрецу Сету. – Назови награду, и ты получишь её.

– Велика доброта твоя, повелитель, – ответил мудрец. – Выдай мне за первую клетку шахматной доски одно пшеничное зерно.

– Одно пшеничное зерно? – изумился царь.

– Да, повелитель. За вторую клетку прикажи выдать два зерна, за третью – четыре, за четвертую – восемь, за пятую – шестнадцать, за шестую – тридцать два...

– Довольно, – с раздражением прервал его царь. – Ты получишь свои зёрна за все 64 клетки доски, согласно твоему желанию: за каждую вдвое больше против предыдущей. Но знай, что просьба твоя недостойна моей щедрости. Ступай! Слуги вынесут тебе твой мешок с пшеницей.

Сета улыбнулся, покинул залу и стал дожидаться последствий своей просьбы у ворот царского дворца. А развязку этой истории вы все, конечно, знаете.

Да, наше мышление, по самой сути своей, линейное – мы привыкли наблюдать постепенное приращение чего бы то ни было: вода в реке или в море поднимается или убывает медленно, растения, животные и даже наши дети растут год за годом и по чуть-чуть, так же незаметно меняются климат, отношения между людьми и т. д., и т. п. Всё постепенно.

Вот почему экспоненциальный рост, о котором говорит Курцвейл в своём законе ускоряющейся отдачи, для нас контринтуитивен: мы не привыкли так думать, а потому и не можем представить себе его последствий.

Царь Шерам решил, что мудрец Сета попросил у него мешок пшеницы, не больше. Но такова действительная реальность экспоненциального роста: если соблюсти последователь-

² Проект «Геном человека» (The Human Genome Project, HGP) – это научно-исследовательский проект по расшифровке последовательности нуклеотидов, составляющих ДНК (порядка 25 тыс. генов). Проект был начат в 1990 году под руководством нобелевского лауреата Джеймса Уотсона, а также Национальной организации здравоохранения США, и стал одной из крупнейших международных научных коллабораций.

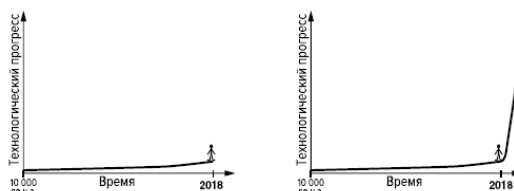
ность, о которой просил мудрец, то к 64-й клетке количество зерна на доске будет в 1800 раз превышать ежегодный современный мировой урожай пшеницы.

То есть Сета попросил у наивного царя весь урожай пшеницы, собранный за всю историю человечества до настоящего момента, а общая масса этого зерна равнялась бы 1200 миллиардам тонн.

В своих прогнозах мы основываемся на опыте тех технологических трансформаций, которые произошли за последние десятилетия. Пережитый опыт диктует нам и наши представления о будущем – таковы особенности мышления человека, его, так скажем, интуиции.

Да, произошёл существенный скачок в развитии технологий – мы все это признаём. «Но что ещё может случиться, чтобы удивить нас? – рассуждаем мы по-стариковски. – Нет, мы уже всё видели...»

Но взгляните на эти два графика: на первом – мы с вами, вместе с тем самым индусским царём Шерамом, а на втором – экспоненциальная кривая технологического прогресса и те самые миллиарды тонн зерна, которые причитаются мудрецу Сете и о которых нас предупреждает Рэй Курцвейл.



Во многом именно благодаря изобретению закона ускоряющейся отдачи Билл Гейтс назвал Рэя Курцвейла «лучшим из тех, кого я знаю, в предсказании будущего искусственного интеллекта».

По оценкам независимых экспертов (уж не знаю, как именно они это измеряли), 86 % прогнозов Рэя Курцвейла «сбывались с высокой точностью».

И действительно, даже если закрыть глаза на эти проценты, прогнозы Курцвейла сбываются как по волшебству – телефоны с bluetooth, синхронный компьютерный перевод, Siri, 3D-видео и очки с дополненной реальностью, суперкомпьютер IBM Watson, машины без водителей и т. д., и т. п.

НИЧЕГО ЛИЧНОГО, ПРОСТО ФАКТЫ

В 1990 году Рэй Курцвейл предсказал, что компьютер победит лучшего игрока по шахматам в 1998 году. Он ошибся: суперкомпьютер Deep Blue компании IBM обыграл Гарри Каспарова на год раньше – в 1997-м.

Тогда же – в 90-м – Курцвейл высказал предположение, что в 2010 году компьютеры смогут отвечать на вопросы, имея беспроводной доступ к информации. Это, как вы понимаете, тоже случилось чуть раньше.

А вот с экзоскелетами, например, великий прогнозист слегка поторопился. Он был уверен, что они позволят инвалидам ходить уже в начале 2000-х, что произошло чуть позже и не повсеместно. Впрочем, соответствующие технологии действительно созданы и активно используются (в частности, компанией Ekso Bionics).

Спустя десять лет – на пороге нынешнего тысячелетия – Курцвейл тоже сделал несколько чрезвычайно смелых прогнозов. Так, например, он обещал, что к 2009 году компьютер будет воспринимать голосовые команды. Случилось это не в 2009-м, но кто из нас не общался с Siri, OK Google или Алисой?

В том же 2009 году Курцвейл ожидал появления очков, стекла которых будут оснащены дисплеями, воспроизводящими эффект дополненной реальности. Вроде бы и тут ошибся – прототипы Google Glass появились только в 2011-м. Но и эти экраны, и технология дополненной реальности появились даже до 2009 года. Так что всё ОК.

Ещё через пять лет – в 2005 году – Курцвейл предсказал, что к 2010 году появится возможность осуществлять языковые переводы с одного языка на другой в режиме реального времени. Skype Translate Microsoft, Google Translate и другие технологии справились с этой задачей. Некоторые же приложения, как, например, Word Lens, и вовсе могут переводить слова на изображении с вашей камеры.

Помню, когда ко мне в гости в интеллектуальный кластер «Игры разума» приехал главный художник Google, автор культового романа «Поколение X» Дуглас Коупленд, с которым у нас перед этим состоялась заочная дискуссия о будущем искусственного интеллекта, это приложение только вышло. И они с куратором его выставки Марселлом Дантесем как малые дети бегали по нашим зданиям, прикладывая свои iPhone к указателям, и радовались эффектам – на экране то же видеоизображение, а текст меняется на английский.

В 2010 году Курцвейл обещал, что к 2019-ому «провода и кабели для персональных и периферийных устройств любой сферы уйдут в прошлое». Что ж, взгляните на наушники к iPhoneX, беспроводные зарядные устройства для Samsung Galaxy S6 (Wireless Charging Pad) или Cota Wireless Power – универсальную колонку для зарядки электроприборов с диаметром действия больше 10 метров.

Вам не кажется, что Курцвейл даже как-то запаздывает со своими прогнозами?..

Кривая закона ускоряющейся отдачи Рэя Курцвейла предполагает наличие трёх последовательных фаз:

- первая – медленный рост (ранняя фаза экспоненциального роста);
- вторая – быстрый рост (взрывная фаза, когда кривая стремительно взмывает вверх);
- третья – фаза стабилизации, когда формируется принципиально новая технологическая парадигма.

И если кому-то кажется, что стабилизация уже наступила, – не обольщайтесь. Вот что обещает нам Рэй Курцвейл на ближайшие десятилетия...

Считается, что вычислительная мощность нашего мозга равняется примерно десяти терабайтам – это очень-очень много, и стоимость такого компьютера сейчас, как вы понимаете, почти космическая.

Но посмотрите на свой телефон: аналогичную вычислительную мощность в 60-х вы могли бы купить лишь за триллион долларов, а в начале 80-х прошлого века – за миллиарды долларов. Но вряд ли сейчас ваш телефон стоит дороже тысячи, правда? Впрочем, его начинка, поверьте, куда дешевле – вы переплачиваете за программное обеспечение и бренд.

Теперь внимание: по расчётам Курцвейла, десять терабайтов – мощность, равная мощности нашего мозга, – обойдутся нам в 2020 году всего в одну тысячу долларов. Проще говоря, к этому моменту персональные компьютеры не только достигнут вычислительной мощности, сравнимой с человеческим мозгом, но будут общедоступны.

В 2011 году журнал Science опубликовал статью Мартина Хильберта из Университета Южной Калифорнии, где он писал следующее: «Люди всего мира могут осуществить $6,4 \times 10^{18}$ операций в секунду на обычных компьютерах образца 2007 года, что сравнимо с максимальным количеством нервных импульсов, возникающих в одном человеческом мозге за секунду».

Самый быстрый суперкомпьютер в этом же 2011 году обладал мощностью 10,51 петафлопс (10,5 квадриллионов операций в секунду – то есть как минимум на две степени меньше, чем требуется для воспроизводства мощности, соответствующей человеческому мозгу).

В 2013 году группе немецких и японских исследователей удалось симулировать одну секунду активности одного процента мозга человека, правда за 40 минут и на кластере из 82 944 процессоров.

В 2018 году был представлен американский суперкомпьютер Summit, производительность которого, по заверениям создателей, приближается к 3,3 эксаопсам, а это три с лишним квинтиллиона операций в секунду – то есть те самые «миллиарды миллиардов», о которых говорил Мартин Хилберт, рассчитывая мощность человеческого мозга.

Теперь представим, что мы перешагнули 2020 год и вошли в будущее, стоящее у нас на пороге. Итак, чем же ознаменуются для нас ближайшие десятилетия?

СПРАВОЧНО

Увеличение мощности компьютеров, уменьшение их размеров и снижение цены производства – следствия так называемого «закона Мура», который был сформулирован больше пятидесяти лет назад Гордоном Муром – тогда ещё только будущим основателем компании Intel.

На самом деле это просто эмпирическое, то есть основанное на опыте наблюдение: мощность компьютеров, обусловленная увеличением количества транзисторов, уместяющихся на кристалле интегральной платы, а также ростом их таковой частоты, удваивается каждые 18 месяцев. Тогда как с ценой происходит обратная ситуация – примерно каждые два года она в два раза уменьшается.

Журнал «Scientific American» привёл такую аналогию: если бы авиапромышленность последние 25 лет следовала «закону Мура», то сейчас Boeing 767 стоил бы 500 долларов, совершал облёт земного шара за 20 минут и затрачивал на это менее 20 литров топлива.

Да, цифровой мир живёт по своим законам – Мура и Курцвейла.

Повсеместная роботизация

Сейчас даже растёт количество людей, считающих меня слишком консервативным и сдержанным в своих прогнозах.

РЭЙ КУРЦВЕЙЛ

С 2020 по 2030 год нас ждёт повсеместная роботизация.

В США и Европе будут приняты первые законы, регулирующие отношения людей и роботов. Деятельность роботов, их права и обязанности будут превращаться в своего рода «Кодекс поведения робота». Впрочем, предполагается, что этот кодекс будет налагать определённую ответственность и на пользователей – то есть на нас с вами.

Это кажется почти что абсурдным. Но посудите сами...

В 2018 году на международной выставке высоких технологий Computex был представлен чип под названием Jetson Xavier (ещё его зовут Исаас). Компания-производитель NVIDIA утверждает, что это идеальный мозг для роботов, ориентированных на использование искусственного интеллекта и глубокое машинное обучение. На его создание были затрачены около 8000 человеко-лет работы и куча денег.

«При энергопотреблении в 30 ватт Jetson Xavier имеет практически такую же вычислительную мощность, как и огромные рабочие станции стоимостью в 10 000 долларов, но при этом он стоит гораздо дешевле», – говорит президент NVIDIA Жэньсюнь Хуан. Кстати, его надо поздравить: Исаас и правда выходит на рынок и стоит чуть больше тысячи долларов.

Впрочем, уже в середине третьего десятилетия нашего века и сами люди подвергнутся киборгизации – Рэй Курцвейл ожидает появление массового рынка гаджетов-имплантатов.

Да, киборг – уже не вымысел. Не буду рассказывать про коленные протезы, оснащённые самообучающимся искусственным интеллектом («RheoKnee» компании Ossur), – им уже больше десяти лет, и это скучно. Куда важнее то, что можно делать с нашим мозгом.

После того как Уильям Доббел создал в 2002 году первую технологию, позволяющую переводить изображение с обычной видеокамеры непосредственно в мозг человека (минуя глаза, глазные нервы и прочие анатомические «излишки»), мы оказались в поистине новой реальности.

Теперь совершенно очевидно, что наш хваленый мозг – это просто серверное пространство. И тот же Рэй Курцвейл уже работает над новым для мозга программным обеспечением, а также над нестандартными средствами доставки в него информации.

У наших органов чувств масса ограничений, но почему бы не подключить к нашему мозгу, например, электронный микроскоп или супермощный телескоп? Почему бы не перепрограммировать наш мозг, зная его «язык», его «код» и весь набор «уязвимостей нулевого уровня»?

Биологический и интеллектуальный апгрейд – это не вымысел фантастов, а уже почти наступивший дивный новый мир.

Параллельно с нашей собственной киборге-низацией, по прогнозам Курцвейла, и персональный робот, способный на сложные, полностью автономные действия, станет настолько же привычной вещью в наших домах, как и бытовая кухонная техника.

К этому Курцвейл добавляет беспроводной доступ к интернету, который покрывает практически всю поверхность Земли, а также солнечную энергию – настолько дешёвую и доступную, что человечеству больше не потребуется углеводородное топливо.³

СПРАВОЧНО

Курцвейл делает ставку на естественные источники энергии, где действительно достигаются всё новые и новые, поражающие воображение успехи. Но возможен и другой способ решения энергетической проблемы – например, термоядерный синтез¹.

Для того чтобы реакция термоядерного синтеза пошла, необходима разогретая до невероятной температуры плазма. Этот эффект уже достигается в специальных установках, которые называются токамак. Нагрев и удержание плазмы осуществляются в них с помощью магнитного поля огромной силы. Сейчас учёным удаётся разогревать плазму до 50 миллионов градусов по Цельсию и продержат в таком состоянии больше 100 секунд (рекорд принадлежит китайской установке EAST).

Это много с точки зрения достигнутого прогресса и недостаточно для полного решения задачи. Однако специалисты уверены, что проблема будет решена уже в ближайшие годы.

К середине 2020-х, обещает Рэй Курцвейл, будет побеждено большинство болезней. Всё это благодаря нанороботам, которые продемонстрируют способность справляться с болезнями лучше любых современных медицинских технологий и лекарств.

Мини-компьютеры величиной меньше, чем клетка крови, будут перемещаться в наших телах и, связанные с облаком, сначала предупредят «Большого Доктора» о рисках наступления заболевания, а затем, если потребуется, успешно его излечат.

Один из основоположников наномедицины Роберт Фрейтас говорит: «Нанотехнологии уже производят биороботов. Примерно к 2020 году возникнут гибридные роботы на основе усовершенствованной ДНК, синтетических белков и других небиологических материалов. В начале 2030 года или раньше учёные построят полностью искусственные устройства: нанороботы, управляемые компьютерным программным обеспечением и способные защитить каждую клетку организма от болезней и травм».

В 2017 году команда учёных во главе с Джеймсом Туром из Университета Райса опубликовала статью в журнале *Nature*, где была представлена, например, нанотехнология, «взламывающая» и «высверливающая» мембраны раковых клеток.

Структура этого молекулярного комплекса действительно повторяет обыкновенную электродрель: основа – это неподвижный статор, который закрепляется на мембране «распорками», воздействие ультрафиолетовых лучей запускает ротор, и тот начинает вращаться со скоростью 2–3 млн оборотов в секунду, буквально пробуравливая мембраны клеток и оставляя открытые поры (см. рис. № 1).

³ Термоядерный синтез – получение более тяжёлых ядер элементов с выделением огромного количества дешёвой энергии.

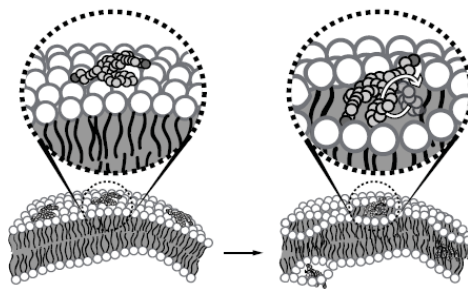


Рисунок № 1. Схема работы наноробота команды Дж. Тура

Учёные могут модифицировать функциональные группы на статоре наномашин так, чтобы они крепились только к определённым структурам клеточных мембран. Это позволяет делать их действия избирательными и атаковать только больные клетки. Пока подобные фокусы проделывают на мышах, но уже тот факт, что мы можем вылечить рак простаты у мыши, не повредив здоровые ткани, – это безусловная победа нанотехнологий.

Фрейтас вместе с Курцвейлом считают, что старение и «естественная смерть» являются *заболеваниями*, которые возникают, когда клеточная структура организма не может восстановить нанесённый ей ущерб. Если лечить эту «болезнь», то молодые смогут оставаться молодыми, а старые помолодеют.

Да, умереть в новом дивном мире будет всё сложнее и сложнее...

Больше тридцати пяти тысяч человек ежегодно гибнут в России в результате дорожно-транспортных происшествий, чуть меньше – в США, примерно такие же показатели – в объединённой Европе. И основная причина – человеческий фактор, то есть водители, которые, например, засыпают за рулём или садятся за него в изрядном подпитии.

Так что не удивляйтесь: во избежание человеческих жертв нам скоро запретят выезжать на дорогу самостоятельно – только в автомобилях с автопилотом.

Начнётся, впрочем, всё с малого – в начале 20-х под запрет попадут автомобили, не оборудованные компьютерными помощниками, а затем и вовсе – только автопилот, и никакой вам дорожной импровизации.

Когда же окажется, что водить автомобили нельзя, то и желающих содержать собственное авто, которое в среднем 98 % времени стоит на приколе, будет глупо и абсолютно невыгодно.

Городские магистрали наполнятся самоуправляемыми автомобилями-такси и, возможно, даже сузятся из-за тотальной уберизации в пользовании автотранспортом (по расчётам Курцвейла – это где-то 2033 год).

Такси-беспилотники будут постоянно разъезжать по вызовам, а парковки в нынешнем их виде просто исчезнут. Кроме того, машины, объединённые общей сетью, начнут договариваться между собой, чтобы избежать пробок.

Эти трансформации будут происходить настолько стремительно, что мы и не заметим того, как нынешние «автолюбители» исчезнут как класс. К хорошему, как известно, быстро привыкаешь. Уже сейчас молодые люди, пользующиеся относительно дешёвым такси, не спешат обзаводиться водительскими правами, а скоро этого вообще не нужно будет делать.

Осознавая это, автоконцерны, выпускающие традиционные автомобили, агрессивно и наперегонки инвестируют в автопилоты.

Да, возможно, Tesla сейчас единственный стопроцентный беспилотник. Но не думайте, что это странная причуда наивного визионера-фантаста Илона Маска. Автомобили, которые вы уже сегодня покупаете у мировых автогигантов, имеют усовершенствования, призванные сделать их совместимыми с автопилотом.

Всё, что через паузу останется автопроизводителю, – это доустановить на вашу машину сканирующий лидар (специальное устройство, обеспечивающее получение и обработку оптической информации) и подключить его к системе искусственного интеллекта, управляющего автотранспортным средством (см. рис. № 2).

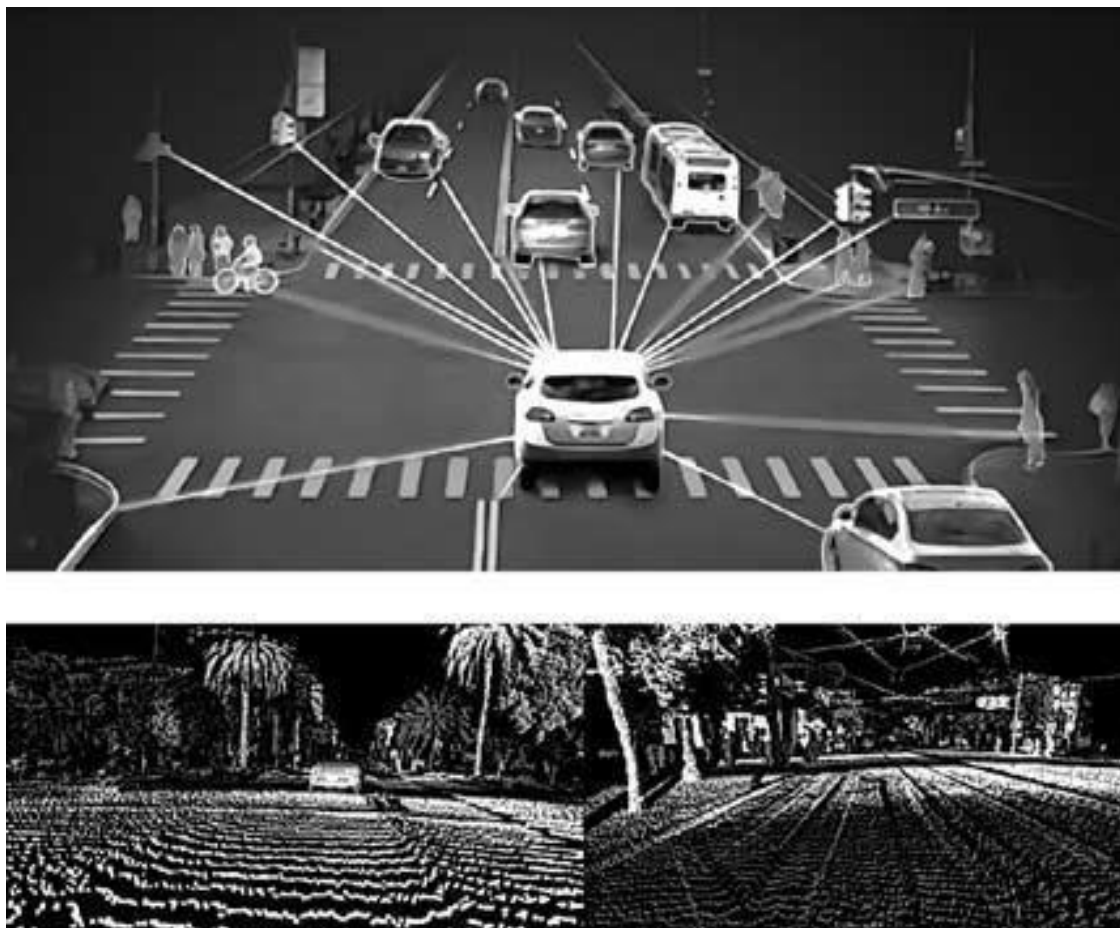


Рисунок № 2. Схематичное изображение того, как лидар реконструирует окружающую его среду

В 2012 году, когда Google только тестировал первые лидары, их цена составляла порядка 70 тысяч долларов за штуку. Уже через два года цена упала до тысячи долларов, а последнее поколение стоит порядка 90 долларов за экземпляр, который вдобавок и размером-то теперь не многим больше почтовой марки.

Впрочем, зачем все эти сложности – не вполне понятно...

Новое соглашение между Uber и NASA предполагает создание летающих такси с вертикальным взлётом. Промо-ролик этого проекта выглядит как научно-фантастический фильм: героиня, попавшая в дорожную пробку, опаздывает на семейный ужин, но на крыше ближайшего небоскрёба её ждёт аэротакси, которое с комфортом доставит её домой.

Кстати, мэр Лос-Анджелеса Эрик Гарсетти уже пообещал проекту свою помощь и считает, что его город представляет собой идеальный полигон для его реализации.

«Но в какую же баснословную сумму обойдётся подобная поездка?» – спросите вы. Создатели Uber Elevate утверждают, что расценки окажутся такими, что горожанам будет дороже содержать собственный автомобиль. Думаю, что в булочную на таком устройстве вы вряд ли поедете, но расстояния побольше будет комфортнее преодолевать именно таким образом.

Если, конечно, вам это вообще понадобится...

ДЕЦЕНТРАЛИЗАЦИЯ

Если верить Рэю Курцвейлу, то мы движемся к тотальной децентрализации всего и вся.

В скором времени человек уже не будет привязан к рабочему месту – на производстве его заменят роботы, а офисы и вовсе исчезнут за ненадобностью. Сквозной блокчейн сделает все наши взаимодействия предельно прозрачными, так что работодателям и заказчикам уже не нужно будет никого контролировать.

Системы коммуникации достигнут такого уровня, что вы будете испытывать ощущение физического присутствия собеседника. А если так, то зачем вам куда-то ехать, чтобы с ним встретиться?..

Архитектура развлечений, в свою очередь, не потребует больше ни театров, ни кинотеатров, ни аттракционов, ни стадионов, ни даже путешествий – все соответствующие впечатления придут к нам в дом сами. Так что жить в больших городах будет просто бессмысленно, с учётом экологии, шума, криминала, эпидемиологической нагрузки.

В общем, вполне вероятно, что нам предстоит беспрецедентная деурбанизация. А оказавшись за городом, но имея всё, что нужно для работы, общения и развлечений, мы даже перемещаться куда-либо не захотим.

Так что да – беспилотники понадобятся, но для экстренных и единичных случаев. Основная же их задача в будущем – это грузоперевозки. Впрочем, и те под вопросом, если вопрос с 3D-печатью решится так, как предсказывает тот же Курцвейл.

По этой же причине и ракеты, которые Илон Маек собирается запускать в качестве нового трансокеанического средства передвижения, возможно, будут простаивать. Да и летающие автомобили Uber Elevate могут не понадобиться⁴.

И раз уж мы заговорили о развлечениях... К 2020 году на рынок поступят специальные очки, которые смогут проецировать любое изображение непосредственно на сетчатку нашего глаза.

Многие помнят скандал вокруг «умных очков» Google Glass – мол, как-то это неприлично стримить в интернет незнакомцев, которые не дали вам на это своего согласия.

Но прогресс, честно говоря, застопорился из-за сугубо технической проблемы: умным очкам, чтобы поддерживать потоковое видео, нужна большая батарея. Но кто будет таскать на себе огромные очки? Они и нос, и уши отдавят. Даже заряда смартфона иногда на сутки не хватает, а он сам по себе не маленький.

Однако закон ускоряющейся отдачи продолжает бить рекорды... Учёные из Университета Вашингтона уже создали технологию потокового вещания, которая потребляет в 10 000 раз меньше энергии, чем прежние. Хитрость в том, чтобы, используя умную беспроводную передачу, выгрузить большую часть работы на другое устройство.

Обычно камерам приходится обрабатывать и сжимать видео перед передачей сигнала по беспроводной сети, а новый подход напрямую прикрепляет пиксели с камеры к беспроводной антенне и отправляет данные импульсно на ближайшее устройство (телефон, планшет или ПК), которое уже и занимается обработкой сигнала.

⁴ Впрочем, это лишь логика рассуждений, которую предлагает нам Рэй Курцвейл. Справедливости ради отмечу, что, согласно знаменитому прогнозируемому перечню IBM «5 in 5», к 2030 году в городах развивающихся стран будет проживать более 80 % населения. Кто ошибается, а кто прав – пока непонятно.

Первый прототип ограничен десятью кадрами в секунду и небольшим расстоянием, но это только начало. А практическое применение этого ноу-хау уже очевидно: можно будет носить умные очки с маленькими батареями или даже вообще обходиться без них (если, например, они будут получать энергию от радиосигналов, что тоже возможно).

Теперь представьте себе спортивное соревнование, за которым вы наблюдаете от первого лица – видите, как ударяете по мячу в финальном матче Чемпионата мира по футболу, попадаете в мишень на Олимпийских играх, вырывая тем самым победу у соперника по биатлону, мчитесь по трассе Формулы 1...

Вы всё ещё заходите толкаться в задних рядах? Что-то я сомневаюсь. Боюсь, ваш запрос будет выглядеть иначе – например: а можно всю эту радость прямо мне в мозг? Можно – к 2030 году, как обещает Курцвейл, виртуальная реальность станет неотличимой по субъективному восприятию от реальности физической.

ОЖИВШИЕ ВЕЩИ

Впрочем, прежде чем виртуальная реальность окончательно поглотит нас, мы, вероятно, ещё увидим (если, правда, заметим) мир оживших вещей.

Термин «интернет вещей» возник совсем недавно – в 1999 году – и принадлежит авторству Кевина Эштона, который являлся одним из основателей Центра Auto-ID при Массачусетском технологическом институте, где разрабатывались системы подключения объектов физического мира к беспроводным каналам связи.

Итогом этой работы учёных стал так называемый «электронный код продукта» (универсальная система идентификации потребительских товаров), впервые опробованный в логистике компании Procter&Gamble, а теперь завоевавший – в разных видах и формах – весь остальной мир.

Так что, когда вы покупаете обычное молоко в супермаркете, а кассир считывает штрих-код на его упаковке, вы, сами того не зная, принимаете участие в большой-пребольшой игре, придуманной создателем термина «интернет вещей». Впрочем, эта игра только началась.

По оценкам экспертов, мировой рынок умных городских систем к 2020 году достигнет 400 млрд долларов, а рынок интернета вещей перевалит за 0,5 трлн. В 2025 году компании будут зарабатывать на интернете вещей 14,4 трлн долларов в год.

Мусорные бачки уже сейчас способны оповещать коммунальные службы о том, что их пора вычистить, системы геолокации способны регулировать городской трафик, а эпидемиологи могут практически в реальном времени отслеживать распространение вирусов.

То есть речь идёт не просто об «умных холодильниках» и «мыслящих микроволновках»: это прежде всего инфраструктурные вещи – накопители электроэнергии, как те, что уже создали в Tesla, умные термостаты, умное освещение с датчиками, реагирующими на присутствие человека, и т. д., и т. п.

«Умные холодильники», которые сами заказывают необходимую вам еду, отслеживают состояние хранящихся в них продуктов и дают рекомендации по питанию, – уже не новость: Samsung, например, активно их продаёт.

Смартфоны тем временем обучаются распознавать запахи и вкусы: компания Adamant Technologies из Сан-Франциско предлагает технологию в пять раз более чувствительную, чем наш с вами рецепторный аппарат. Так что и свежесть продуктов, и причины запаха изо рта скоро будут определяться не гипотетически, а предметно и в реальном времени.

Впрочем, пока эти холодильники и датчики – лишь первые ласточки. Существенные изменения в качестве нашей жизни произойдут, когда накопится критическая масса таких устройств, датчиков и связанных с ними систем. Если телефон есть только у вас, толку в нём немного, но когда каждый человек постоянно держит его при себе – это другое дело.

По оценкам компании Cisco Systems, занимающейся созданием интернет-решений для бизнеса, к 2020 году количество подключённых к интернету вещей будет превышать 50 миллиардов, а ещё через паузу все 1,5 триллиона «вещей» нашего с вами мира будут объединены в одну большую сеть.

С другой стороны, инновации всегда копятся годами (прежде, впрочем, на это уходили десятилетия или даже столетия), а потом выстреливают всей своей накопленной мощностью уже в совершенно другом, новом качестве.

В конечном итоге нас ждёт «интернет всего» (этот термин изобрели сотрудники той самой компании Cisco Systems) – объединение «интернета вещей» с «интернетом людей».

Пока ещё люди находятся в рамках своей цифровой зоны, лишь сообщая вещам о том, что они делают, что им важно, чего они хотят и т. д. Но скоро «вещи» будут настолько умны, что они сами выйдут с нами на контакт – они будут сами понимать, чего мы хотим, что для нас важно, и совершат все необходимые действия, чтобы мы остались довольны.

Утром, когда вы проснётесь и откроете глаза, умный дом поприветствует вас и сам раздвинет шторы, сделает тёплым пол на вашем пути в ванную, включит воду, когда вы окажетесь перед раковиной, и порадует завтраком из ваших любимых и обязательно свежих продуктов, когда вы дойдёте, наконец, до кухни. Впрочем, это только начало дня...

«Устройства станут самостоятельно общаться, у них появятся свои “социальные сети”, которые они будут использовать для обмена и накопления информации, а также автоматического управления и активации. Мало-помалу мир людей станет местом, где решения принимаются активным набором взаимодействующих устройств. Интернет станет более распространённым, но менее осязаемым, менее видимым. В некотором смысле он станет фоном для всего, что мы делаем».

ДЭВИД КЛАРК, Массачусетский технологический институт

Взбесившийся 3D-принтер

Наше будущее просто усилит то, чем мы являемся на сегодняшний день.

РЭЙ КУРЦВЕЙЛ

С 2030 года по 2040-й нас ждёт замена реального виртуальным, а из реального, кажется, останется только 3D-принтер.

Развитие технологии 3D-печати уже идёт полным ходом, а после 2030 года она, согласно прогнозам Рэя Курцвейла, совершит настоящий переворот в экономике.

Пока технология 3D-печати только ищет себя, хотя она уже неплохо справляется с созданием медицинских протезов и уверенно строит даже многоквартирные дома. Впрочем, бытовые 3D-принтеры пока на какие-то невероятные чудеса не способны.

Но, как говорится, лиха беда начало... В 20-х годах, по мнению Курцвейла, они станут для нас привычной бытовой техникой, а к 2030-ому мы, например, будем печатать с их помощью одежду (интернет будет полон бесплатными моделями для печати – скачивай, печатай и одевайся).

Причём нам даже дизайнеры для создания новых моделей одежды не понадобятся – уже сейчас искусственный интеллект преуспел в разработке дизайна одежды, предметов интерьера и других вещей.

Например, компания Stitch Fix использует программное обеспечение, которое учитывает индивидуальные физические параметры человека, его предпочтения и модные тренды. Комбинируя понравившиеся клиенту варианты вырезов, рукавов и других элементов одежды, оно работает как портной модного дома – создаёт удобную и модную одежду, но, по сути, по вашему индивидуальному заказу.

Если же модные бренды ещё всё-таки сохраняют свою символическую притягательность (что не факт), то защищённые авторским правом цифровые модели одежды от каких-нибудь Gucci и Dsquared мы сможем купить на специальных облачных стоках. А печатать уже, соответственно, «самопалом».

В настоящий момент 3D-печать только пробует себя на этом рынке. Хороший пример – американский стартап Feetz, который занимается производством обуви. Проводятся и первые эксперименты по печати дизайнерской одежды, и даже в России.

Конечно, это пока лишь прототипы, а не отработанная технология, и удобство у этих изделий сомнительное. Но вспомните свой первый кнопочный телефон... А смартфонам ведь немногим более десяти лет.

«Вскоре большинство вещей будет результатом информационных технологий, включая одежду, которую будут печатать на 3D-принтере. С помощью вертикального сельского хозяйства мы сможем выращивать продукты и печатать их на 3D-принтере. К 2020-м годам 3D-дизайн будет настолько доступен, что жить станет намного легче, и мы сможем напечатать всё, что нам будет необходимо, включая дома. [...] Наше общество построено так, чтобы для существования нужна была работа, но всё будет иначе. У нас будут способы предоставить высокие условия жизни для каждого в течение 15–20 лет».

РЭЙ КУРЦВЕЙЛ (2015 ГОД)

К 2030 году 3D-принтеры станут пищевыми, то есть с их помощью мы будем готовить продукты. Точнее – создавать и готовить.

Прототип пищевого 3D-принтера уже работает – его создали по заказу НАСА специалисты компании Systems & Materials Research Corporation. Он смешивает около 12 питательных компонентов, выдавая вполне съедобную кашу.

И это, опять-таки, только начало, а дальше, как предполагается, будут созданы гигантские морские фермы, где станут выращивать дешёвую и питательную еду. Роль «чернил» в наших пищевых 3D-принтерах будут выполнять переработанные особым образом водоросли и морепродукты.

Если же вам всё-таки почему-то захочется настоящего мяса, то его изготовят из стволовых клеток и запихнут всё в тот же принтер.

Вам останется только выбрать на экране компьютера блюдо, которое вы хотите, и это устройство само всё сделает – милости просим, приятного аппетита!

Всё это счастье ожидается до 2030 года, а то, что произойдёт после, – и вовсе за гранью фантастики. По прогнозам Курцвейла, в 2031 году 3D-принтеры будут стоять во всех больницах, а печатать на них будут не только лекарства и инструменты, но и человеческие органы для пересадки.

Вкупе с медицинскими нанороботами это, по идее, должно решить все возможные медицинские проблемы.

СВЕРХЧЕЛОВЕК

Работы по созданию сверхчеловека идут полным ходом.

Мы привыкли думать, что эволюция – это такая дряхлая, медлительная тётенька, которая никак не может перейти дорогу. Мол, мутации – это дело случая, а если какие-то скачки и происходят, то лишь по причине глобальных экологических катастроф.

Но современные технологии в области медицины, биологии, генной инженерии и электроники – стрит-рейсеры. Они рассекают по дорогам мироздания, плюя на правила, и уже давно сбили Тортиллу биологической эволюции аккурат на том самом пешеходном переходе.

Давайте взглянем в лицо фактам: за двести тысяч лет существования человека на этой грешной земле его средняя продолжительность жизни выросла примерно в два раза. Это, конечно, большой рывок. Но только за последние сто лет она выросла ещё в два раза. Сравните: 200 000 и 100.

Согласитесь что-то явно пошло не так. И продолжает идти. Причём возраст – это только один из показателей. Есть ещё рост, вес, патоморфоз болезней, уровень /O...

Тот путь, по которому развивается современное человечество, одни учёные гордо называют «геннокультурной эволюцией» (Кевин Лаланд), другие – «технико-физиологической эволюцией» (Роберт Фогель), третьи – уже без особого энтузиазма – эволюцией «метабиологической» (Джонас Солк). И, пожалуйста, не думайте, что это бред. Пусть само по себе это ничего и не значит, но Фогель и Солк – это всё-таки нобелевские лауреаты.

Мы радикально недооцениваем то воздействие, которое цивилизация оказывает на наш биологический вид. **В руках человечества теперь невероятный потенциал реконструкции нашей с вами базовой биологической матрицы. И возможности, которые здесь открываются, практически безграничны.**

После того как Крейг Вентер синтезировал в 2010 году первую в мире искусственную бактерию, а проще говоря – *искусственную жизнь*, трудно представить, что хоть кто-то всю эту вакханалию сможет

остановить. «Синтетическая биология» уже обещает нам возможность полного клонирования нашей ДНК с целью её последующего ремонта.

И американский Конгресс может сколь угодно торжественно запрещать клонирование и работу со стволовыми клетками, но в чём смысл? Учёные просто переезжают в Китай, Израиль, Японию, Австралию или Сингапур, где их ждут с распростёртыми объятиями и новейшими лабораториями.

С 2030 по 2037 год, во многом благодаря тем же нанороботам, произойдёт фундаментальный прорыв в понимании принципов работы человеческого мозга, а это открывает путь и к цифровому бессмертию, к слиянию нас с суперкомпьютерами и... к полной новых впечатлений сексуальной жизни.

Первое свидание с человекоподобным искусственным интеллектом, по заверениям Курцвейла, ждёт нас уже в 2034 году. Причём свидание в прямом смысле этого слова...

СЕКСУАЛЬНЫЕ ПЕРСПЕКТИВЫ

Всеобщее сексуальное счастье не за горами. Причём каждый сможет встречаться с идеальным партнёром желаемого пола, возраста, внешности, ума, эмоционального склада, темперамента.

Немного он, конечно, будет искусственным, но вы этого не заметите – вас ждёт удовлетворение любых сексуальных фантазий, о которых вы могли только мечтать.

Речь не идёт о забавах с секс-куклой, оснащённой интеллектуальной начинкой, поскольку они и сейчас уже доступны практически в полном объёме, и есть даже специальные бордели, которые предоставляют соответствующие услуги.

Порноиндустрия всегда в авангарде любой технологии. Сообразительные японцы, например, создали уже полноценный «симулятор секса». Он состоит из шлема виртуальной реальности Oculus Rift, специального облегающего костюма, который стимулирует все необходимые части тела, а также искусственной женской груди с обратной связью и навороченного мастурбатора от компании Tenga.

Комплекты первой партии, которые продавались по 400 долларов за штуку, разошлись тут же влёт. Конкретно в данной имитации мужчина становится участником популярной эротической игры SexyBeach, но с учётом того, что порностудии уже активно предлагают виар-видео, то от игры, я думаю, можно легко перейти и к «реальным» персонажам.

Хотя японцам, возможно, секс с героями манги даже интереснее, чем с порнозвёздами. В конце концов, как ещё вы сможете заняться абсолютно реалистичным сексом с мультяшным персонажем?

Да, это новая сексуальная реальность. Поэтому не стоит удивляться, что авторитетный научный журнал Sexual and Relationship Therapy уже опубликовал исследование, обосновывающее введение в номенклатуру новой сексуальной ориентации (идентичности) – цифросексуалы.

Цифросексуалами теперь принято называть людей, которые нуждаются в сексуальном взаимодействии посредством цифровых технологий и не испытывают желания заниматься «обычным» сексом с «живыми людьми». Ввести этот термин сексологов заставило распространение так называемых иммерсивных технологий, обеспечивающих полный эффект присутствия.

Один из авторов этой публикации – доцент Манитобского университета в Канаде Нил Макартур – говорит, что цифросексуалов отличают самые разнообразные специфические сексуальные предпочтения – от непосредственного взаимодействия с роботами (например, с помощью секс-ботов или оснащённых искусственным интеллектом секс-игрушек) до виртуального порно и полного погружения в различные виртуальные среды.

Цифросексуалы часто становятся завсегдатаями специализированных многопользовательских онлайн-игр. Суть проста: вы находите себе пару для «развлечения и досуга», можете вместе походить по виртуальным улицам, барам, клубам, ресторанам, снять номер в гостинице или в специальных местах, оборудованных всем необходимым для специфических сексуальных утех.

Но если цифросексуалу другие реальные пользователи, скрывающиеся за аватарами, не интересны, то можно сойтись просто с роботом. Уже существующие на рынке прототипы секс-роботов неплохо движутся, общаются и, как говорят, очень похожи на людей на ощупь. Специалист в сфере новых технологий из Китая Чжэн Цзяцзя даже, как сообщалось в СМИ, якобы даже женился на созданном им же роботе-женщине Ин-Ин.

Всё это звучит предельно странно, я понимаю. Но, может, это только пока? Журнал *BMJ Sexual & Reproductive Health*, в свою очередь, опубликовал британское исследование, согласно которому 40 % мужчин готовы купить себе секс-бота в ближайшие пять лет, а 49 % мужчин готовы вступить в отношения с «сверхреалистичной» секс-куклой.

Соответствующий бизнес активно развивается. Четыре ведущие в индустрии секс-технологий компании, чья суммарная стоимость оценивается сейчас в 30 миллиардов долларов, продают, прошу прощения за подробности, «реалистичные манекены с переменным возрастом, внешностью и текстурой и настраиваемыми оральными, вагинальными и анальными отверстиями». Стоимость таких секс-роботов достигает 17 тысяч долларов за штуку.

Известный британский футуролог Йен Пирсон уверен, что мы все потихоньку начнём менять «живых» сексуальных партнёров на искусственных уже в течение ближайшего десятилетия. Дело в комфорте и удовольствии: механические машины получают почти совершенный искусственный интеллект, готовый предложить человеку «настраиваемую личность с тем эмоциональным багажом, который он хочет», а потому, говорит Пирсон, секс с роботом «станет проще, безопаснее, чаще и гораздо приятнее».

Прототип такой куклы недавно показали в документальном фильме *The Virtual Reality Virgin* («Девственница виртуальной реальности») на британском телеканале *BBC Three*.

Производитель этой секс-игрушки – Мэтт Макмаллен, фирма которого работает над созданием роботов и программ, помогающих человеку заняться идеальным сексом в виртуальной реальности, – живописует достоинства своего продукта: «Клиент сможет создать из куклы такую личность, какую только пожелает: если захочет, чтобы она была умной, – она будет, не захочет – не будет, если захочет, чтобы она была застенчивой, – кукла будет стесняться». Кроме того, в виртуальной реальности с такой «идеальной девушкой» можно отправиться куда угодно – оказаться, например, в альпийском шале или на берегу океана, или же выбрать что-то более экстремальное.

В фильме есть забавный эпизод, когда ведущий спрашивает виртуальную девушку через специальное приложение: не хочет ли она заняться с ним сексом? Подумав, искусственно-интеллектуальная девушка отвечает: «С удовольствием занялась бы, но пока не могу – я ещё официально не зарегистрирована на владельца. Но я очень люблю трахаться и готова сделать для своего любимого всё что угодно!»

Йен Пирсон считает, что подобные забавы станут в скором будущем не менее распространёнными, чем сейчас обычное порно: к 2030 году тот или иной опыт цифрового секса будет иметь почти каждый половозрелый житель земли, а к 2035-ому, как он говорит, подобные «умные» секс-игрушки обоснуются в большинстве домов.

«Многим втайне нравится идея каким-то образом освободиться от того стесняющего обстоятельства, что для достижения высшей точки сексуального возбуждения требуется другой человек. Но им не стоит обольщаться: последнее слово эротической техники всё ещё выглядит очень неказисто. [...] Секс-робот будущего вполне может быть реализован в виде облегающего костюма со встроенными в него голубыми светодиодами, обращёнными внутрь. [...] Вместо того чтобы задействовать органы чувств естественным путём, как это делал бы секс-робот, нейронная виртуальная реальность будет симулировать этот опыт путём искусственной активации нервных клеток».

ДЭВИД ЛИНДЕН, профессор нейробиологии Университета Джона Хопкинса

Впрочем, согласно прогнозам Курцвейла, к 2038 году мы в принципе будем жить в мире роботизированных людей и сами превратимся в «продукты трансгуманистичных технологий».

Это значит, что вы, ваши друзья и знакомые будут, может быть, оборудованы дополнительным интеллектом – например, ориентированным на конкретную узкую сферу знаний.

Уже сейчас количество информации таково, что ни один профессионал не может знать всего даже в рамках какой-то своей узкой специализации. Это очень показательно, если смотреть на современное состояние медицины.

Если в начале прошлого века врач лечил любую патологию, с которой к нему обращались, то по мере развития медицинской науки оформились отдельные врачебные специальности – кардиолог, пульмонолог, травматолог, невролог, психиатр, абдоминальный хирург и т. д.

В настоящий момент научной информации по каждому отдельному заболеванию столько, что ни один врач не может даже одну болезнь знать досконально. Как тут обеспечишь комплексный подход к пациенту? Приходит на приём не болезнь, а человек целиком.

Получается, что развитие медицины естественным, хотя и парадоксальным образом приводит к снижению качества медицинской помощи: врач, запертый в узких рамках своей специализации, всё больше напоминает слепого мудреца из старой притчи про слона – лечит не то, что надо конкретному пациенту, а то, в чём он разбирается.

Но если соединить мозг человека с цифровым облаком, в котором собран весь объём необходимых данных, и научить его этой базой знаний пользоваться, то уровень компетенции специалиста, конечно, существенно улучшится. Звучит как фантастика, но когда-то и мобильная связь казалась нам чем-то совершенно невозможным и удивительным.

А НУЖЕН ЛИ ВОООЩЕ ДОКТОР?

С другой стороны, я не очень понимаю, зачем нам вообще апгрейдить человека, если можно и вовсе избавить его от этой утомительной необходимости – учиться, работать, стараться что-то делать? Почему бы

просто не убрать его оттуда, где он может быть опасен по причине ограниченности своих возможностей?

И правда, в другой версии будущего максимальной нейтрализации подлежит всякий «человеческий фактор» – не только на транспорте (как я уже рассказывал), но и в той же медицине, например.

По данным корпорации RAND, которая занимается системными исследованиями в этой области, лишь 55 % взрослых пациентов американских медицинских центров получают надлежащую терапию. А это значит, что оставшиеся 45 % прямо или косвенно страдают от врачебных ошибок и некачественно оказанной медицинской помощи.

Разумеется, это бардак, а врачебные ошибки нам не нужны совершенно. Поэтому если люди-врачи уже не справляются с накопленным наукой объёмом знаний, то почему бы их не отстранить от дел вовсе?... И это время очень близко!

Вероятно, вы знаете телепрограмму для умников «Своя игра», но это знаменитая международная франшиза – калька с американского формата «Jeopardy!». В ней у нас побеждали Александр Друзь, Борис Бурда, Анатолий Вассерман. Все как один – интеллектуальные глыбы.

Но с 2011 года любой интеллектуал и эрудит проигрывает в этой игре суперкомпьютеру IBM Watson. Watson не просто знает все ответы (имея свободный доступ к объёмным базам данных, это, наверное, и не слишком сложно), он научился понимать вопросы – различать подтекст, метафору, языковую игру и т. д. Прежде всё казалось невозможным.

Действительно, особенность «Своей игры» («Jeopardy!») в том, что вопросы к участникам формулируются хитро и двусмысленно.

Вот, например, вас спрашивают: «Самопроизвольно закипает и без внешних причин охлаждается, хорошо взаимодействует с металлами одиннадцатой группы таблицы Менделеева, помогает снять стресс и является эффективным чистящим и моющим средством. Что за создание описано в одном журнале?»

Даже человеку непросто догадаться, что это «создание» – женщина. Меня, например, на этот ответ может навести только тот факт, что к металлам одиннадцатой группы относятся золото и серебро. Но как до этого додумывается IBM Watson, непонятно категорически – причём ни нам смертным, ни даже её создателю – Дэвиду Ферручи.

Машина – IBM Watson – просто бьёт вопрос на отдельные элементы, затем обращается к огромной базе данных, получает некие комбинации фактов и сопоставляет одни с другими. Когда же она получает необходимое совпадение – бинго! – жмёт на виртуальную кнопку, а её соперники-люди оказываются посрамлены.

Уже тогда – в далёком 2011 году – мозг Watson'a представлял собой параллельную вычислительную систему, состоящую из девяноста серверов IBM Power 750. Система была способна анализировать 500 гигабайт информации в секунду, или, если перевести на человеческий язык, анализировала примерно 3,6 миллиарда книг в час.

Представьте, какой апгрейд она пережила к настоящему моменту... Тем более что участие IBM Watson в «Jeopardy!» было лишь тренировкой и хорошим рекламным трюком.

На самом деле Watson создан, конечно, не для игр. Он призван полностью заменить врачей и уже неплохо с этой задачей справляется⁵. В конце концов, что такое набор симптомов, которые демонстрирует больной на приёме у врача, как не такой вот каверзный вопрос с набором двусмысленностей?

Как врач скажу, что именно такую задачу и решает ваш доктор, когда изучает кипу анализов и расспрашивает вас о симптомах болезни.

Действительно, ты сначала бьёшь множество фактов на отдельные кластеры, сверяешь их со своим медицинским багажом, получаешь какие-то отдельные вероятностные ответы, уточняешь что-то, чтобы проверить свою гипотезу, а затем смотришь, какое совпадение фактов наилучшим образом укладывается в логику того или иного заболевания.

Только теперь давайте хотя бы гипотетически сравним «багаж медицинских фактов», которые могут находиться в голове одного, пусть даже гениального доктора, с теми базами данных по медицине, экологии, эпидемиологии, фармакологии и т. д., которыми может обладать наш новый добрый доктор Watson... Думаю, даже сравнивать бессмысленно – Watson выигрывает с разгромным счётом!

Да, это поначалу пугает: как – отдать человека и его здоровье на откуп бесчувственной машине?! Но ежегодные инвестиции в обучение Watson'a врачебному искусству исчисляются сейчас миллиардами долларов, и я не преувеличиваю. Так что это просто дело времени. Скоро он будет щёлкать болезни как орехи.

Теперь заглянем ещё чуть-чуть подальше в будущее и представим себе капсулу: вы в неё залезаете, она вас сканирует, делает расчёт с учётом всех существующих на данный момент медицинских знаний и, даже минуя фазу диагноза, выдаёт вам индивидуальное лекарство, учитывающее все особенности вашего организма (включая аллергический статус, например, или непереносимость лактозы). Ни один врач на это не способен и никогда способен не будет. Тогда зачем он вообще нужен? Все в сад, товарищи!

Впрочем, про капсулу я немножко подзагнул – ну хочется хоть чуть-чуть медицинской техники, по старой врачебной памяти! На самом деле капсулу, видимо, заменит обычный смартфон, который будет общаться напрямую с облаком, где заживёт в скором времени тот самый Большой Доктор Watson.

Вы спросите: а как же анализы? Это ещё одна технологическая штука из области научной фантастики, которая, впрочем, уже сейчас в стадии реализации и массового внедрения.

Компания Tribogenics, например, придумала замену рентгеновскому аппарату на основе липкой ленты – я не шучу, что-то именно наподобие липкой ленты из хозяйственных товаров и используют! Универсальной диагностической лабораторией крови скоро станет гидротропный полимер, разработанный Джорджем Уайтсайдсом. Кстати, его можно будет напечатать на обычном к тому времени домашнем 3D-принтере.

Или вот, например, компания Nanobiosym доктора Аниты Гоэл (Анита пока ещё человек) создала «Лабораторию на чипе». Звучит очень помпезно, не правда ли? В действительности это нехитрая нанотехнологическая платформа,

⁵ По иронии судьбы, а может и закономерно, тут снова не обошлось без Курцвейла: над обучением Watson'a медицинским специальностям работает медицинская школа штата Мэриленд, Колумбийский университет и компания Nuance Communications, которая была первым стартапом Курцвейла и называлась тогда Kurzweil Computer Products.

которая по капле вашей слюны или крови определяет ДНК- или РНК-следы любого патогена. Уже сейчас эта технология позволяет за 15 минут по одной-единственной капле крови провести тестирование на ВИЧ. А стоит это удовольствие – закон Мура-Курцвейла! – меньше одного доллара.

Как вы уже, наверное, догадываетесь, все эти анализы, которые вы самостоятельно сможете провести у себя дома, будут автоматически отправлены вашим гаджетом облачному доктору Watson'у. Потом необходимые вам лекарства напечатает, надо полагать, тот же 3D-принтер. И всё, клиники по всему миру можно смело закрывать...

Ещё, правда, осталась хирургия, которую мы здесь не обсудили... Но я пожалею ваше воображение, поскольку даже для меня это уже чересчур. Если же вам всё-таки интересно и не терпится заглянуть в будущее хирургии – погуглите Da Vinci Surgical System. Завораживает, что эта хирургическая система уже делает, и самое главное – чем, как планируется, она будет заниматься в самом ближайшем будущем⁶.

Сейчас, когда врач, например, проводит лапароскопическую операцию, он следит за действиями специальных манипуляторов, находящихся в организме пациента, через экран монитора. А теперь представьте, что он будет «находиться» прямо внутри вашего организма – пусть и виртуально, но с абсолютной достоверностью. Через паузу, впрочем, и самого хирурга уже не потребуется...

«В будущем – через 10~15 лет – мы сможем определять рак по анализу крови, мочи или по дыханию, а определив, удалять опухоль с помощью роботов. Робот найдёт крошечное раковое повреждение, введёт иглу и уничтожит его – точно так же, как мы сейчас уничтожаем злокачественную родинку».

КЭТРИН МОР, руководитель медицинских исследований I Intuitive Surgical

К 2039 году, по заверениям Курцвейла, наномашинки будут имплантироваться непосредственно в человеческий мозг, что позволит нам осуществлять произвольный ввод и вывод сигналов из клеток мозга.

В частности, это позволит нам оказываться в виртуальной реальности с эффектом «полного погружения» без всякого дополнительного оборудования, не считая, разумеется, тех самых нанороботов.

Наконец, в 2040 году поисковые системы, говорит Курцвейл, станут основой для гаджетов, непосредственно вживлённых в человеческий организм. **То есть поиск по информационным массивам мы будем осуществлять уже как бы подсознательно, даже без голосовых команд: вы просто задумались о чём-то, сами того не осознавая, а на экране специальных глазных линз уже прокручивается вся необходимая вам информация.**

Возможно, ещё более шокирующим и противоестественным выглядит прогноз Рэя Курцвейла, согласно которому наши возможности будут расширены не только за счёт прямого доступа к облакам данных, но и посредством разнообразных имплантов – от глаз-камер до дополнительных рук-протезов.

Соответствующие технологии уже активно разрабатываются во множестве научных лабораторий по всему миру. Нейроинтерфейс iBrain, позволяющий мысленно управлять механической рукой или компьютером, как вы, наверное, знаете, активно тестировал на себе Стивен Хокинг.

⁶ Добавлю, что в 2015 году компании Google и Johnson&Johnson объявили о совместном проекте по производству робота-хирурга, который, по их заверениям, будет многократно превосходить возможности Vinci Xi.

Искусственной рукой i-LIMB Pulse, которая позволяет полностью воспроизвести нашу мелкую моторику, пользуется уже более пятнадцати тысяч человек. 200 тысяч человек вернули себе слух благодаря электронным имплантатам. Около тысячи пациентов живут с полностью искусственным сердцем Total Artificial Heart.

А, например, компания EnChroma создала очки, которые позволяют людям, страдающим дальтонизмом и ахроматопсией, раскрашивать мир не только в правильные цвета, но и вообще видеть цвет – то есть перейти от чёрно-белого зрения к цветному.

Сейчас все эти технологии, конечно, очень дороги, но IBM прогнозирует, что начиная с 2022 года подобные увеличивающие возможности человека инструменты будут находиться в серийном производстве.

Так что, учитывая весь этот «трансгуманистический» апгрейд, секс с роботами для таких киборгов, какими мы будем к концу 30-х годов, уже даже сложно будет назвать извращением.

ВОСКРЕШЕНИЕ ПО ЛАЙТУ

Ещё одна из загадочных инициатив Рэя Курцвейла – это воскрешение покойников. Не знаю, как лично вы к этому относитесь, но Курцвейл настроен весьма решительно. Вот как он рассуждает...

Во-первых, каждый из нас оставляет в сети колоссальное количество информации – до 1,5 ГБ в сутки. Это и правда очень много. Особенно если вспомнить, каких результатов удалось добиться исследователю Кембриджской лаборатории психометрии Михалу Косинскому, в чьём распоряжении были лишь единичные пользовательские лайки.

Итак, мы оставляем немислимое количество информации в сети и в какой-то момент умираем (пока это так). Родственники опечалены – все в трауре, страдают и мучаются. Так почему бы не вернуть им нас (пусть и в каком-то урезанном, так сказать, виде)?

Да, план такой: обучить бота быть воскрешённым из мёртвых покойником, который способен общаться со своими родственниками, а возможно, даже и развиваться с учётом того, что происходит в осиротевшем семействе.

Немного шокирует, правда? Но Курцвейл идёт в своих фантазиях куда дальше! Ваши безвременно умершие родственники наследили не только в Сети, но и в вашем собственном мозгу – как бы иначе вы о них тосковали? Так вот, Курцвейл обещает нанороботов, которые смогут считывать воспоминания, хранящиеся в ваших нейронных связях.

Собрав эти воспоминания об умершем непосредственно из вашего мозга, они отправят их в цифровое облако (новое, простите, Царствие Небесное). Там эта информация будет переработана, что позволит воссоздать мертвеца не только по данным из Сети, но и соответствующего вашему собственному ощущению. Отличить подделку от натурального человека в таком случае будет уже невозможно.

Думаю, что подобное воскрешение кажется кому-то странной причудой, болезненной фантазией. Но, опираясь на свой психотерапевтический опыт (а мне много пришлось работать именно с пациентами, переживающими потерю близких), должен вам сказать, что услуга «воскрешения», предлагаемая Рэем Курцвейлом, несмотря на свою «нетипичность», вполне может найти коммерческое применение.

Кстати, как именно это может выглядеть, рассказывает сериал «Чёрное зеркало»: этой фантазии посвящена отдельная серия второго сезона – «Я скоро вернусь». Молодой человек погибает в автокатастрофе, его девушка решается

заказать сначала его виртуальную версию, а потом, чувствуя его живым и настоящим, – в физическом исполнении.

Дальше в страшной фантазии автора сериала Чарли Брукера всё идёт немножко не по плану, и не без драматизма. Но, в конце концов, сериал-то сделан для нас и наших современников, а как люди будут реагировать на подобные предложения лет через двадцать – предсказать сложно.

Так, например (и это уже не фантазия и не далёкое будущее), в 2018 году шведская IT-компания, работающая в сотрудничестве с сетью стокгольмских похоронных бюро «Феникс», сообщила о скором появлении продукта, который позволит нам общаться с цифровыми копиями наших покойных друзей и родственников.

На первом этапе у ещё живых ещё добровольцев записываются голоса, а после их смерти специальный бот обучается поддерживать с вами беседу от их имени. Новое кино вы, я полагаю, вряд ли сможете с ним обсудить, но вот поболтать о домашних делах – вполне.

Следующий этап – это создание визуальной копии мертвеца. Этим также активно занимаются: южно-корейская компания Elrois уже создала мобильное приложение (With Me), с помощью которого вы можете пообщаться с 3D-аватаром друга и даже сделать с ним совместное селфи. Жив этот друг или мёртв, для данного приложения, как вы понимаете, значения не имеет.

В общем, шаг за шагом, и, вполне вероятно, мёртвые вернуться к нам...

Смерть – это, как вы уже, наверное, поняли, панический кошмар гениального изобретателя. Курцвейл называет её «великим похитителем отношений, знаний и смысла», и кажется, всё, что он делает, служит этой – одной-единственной цели – обрести бессмертие.

«Если вы будете хорошо себя чувствовать до 2030 года, вы вполне сможете жить так долго, как вам захочется», – обещает Рэй Курцвейл, которого то ли в шутку, то ли всерьёз называют не только «пророком техноконца», но ещё и «гением техноспасения».

Совместно со знаменитым американским врачом Терри Гроссманом он даже создал компанию, которая предлагает своим состоятельным клиентам продлевать жизнь за счёт нормализации рациона питания, режима дня и внедрения различных технологических новинок.

Рацион самого Курцвейла включает в себя около двухсот пищевых добавок в сутки. А годовая порция этого «здоровья» стоит в районе миллиона долларов⁷.

Цель Курцвейла проста – дожить до технологического бессмертия. По его расчётам, остаётся ещё пару десятилетий, может, чуть больше, а Рэю пошёл уже восьмой десяток... Нужно как-то смочь дотянуть.

Пока главные надежды на существенное продление жизни возлагаются на терапию стволовыми клетками, генную инженерию, 3D-печать органов и другие медицинские процедуры.

Но Курцвейл смотрит дальше, его интересует подлинное бессмертие, а тут без загрузки сознания на цифровые носители – никак.

Солнечная система, в любом случае, в какой-то момент погибнет, а в физических телах покинуть её пределы крайне сложно. Но будучи простым цифровым кодом – вполне.

Энтузиазм Рэя Курцвейла поддерживают и некоторые отечественные учёные, которые создали «Стратегическое общественное движение “Россия-2045”». Как следует из названия, основатель этой инициативы Дмитрий Ицков и его соратники всерьёз рассчитывают, что технологическая сингулярность ждёт нас к 2045 году.

⁷ Не буду, впрочем, утруждать вас деталями: что, почему и зачем, Рэй Курцвейл подробно рассказывает в книге «Transcend: девять шагов на пути к вечной жизни», которая написана им в соавторстве с тем самым профессором Гроссманом.

Главный проект этого общественного движения называется «Аватар», и будет он реализовываться в четыре этапа.

«Аватар А» планируется создать уже к 2020 году. Это будет искусственное тело человека, которое может дистанционно управляться через «мозг-компьютер». В рамках этого этапа планируется создать усовершенствованные протезы органов тела человека и органов чувств, экзоскелеты различных функций и новые человеко-компьютерные языки.

Для «Аватара Б» срок установлен до 2025 года. Тут суть в том, что учёные создадут искусственное тело, в которое можно будет произвести трансплантацию мозга человека, биологический организм которого уже израсходовал свой ресурс – в общем, перед смертью. Что делать со старением самого мозга, пока непонятно, но это дело, как считается, наживное – омолодят.

Кроме того, в планах также «Аватар В», где в искусственное тело будет перенесена «нематериальная структура сознания» (что бы это ни значило...). Предполагается, что это возможно благодаря «обратному конструированию» мозга. Ждём новинку к 2030–2035 годам.

Наконец, аккурат к технологической сингулярности должен появиться «Аватар Г» – тело из нанороботов и тело-голограмма. Звучит как абсолютное безумие, но о том же говорит и Рэй Курцвейл: этот его прогноз и изображение тела, созданного из облака нанороботов, уже попали на обложку «Time».

Комиссия по борьбе с лженаукой РАН, впрочем, обеспокоилась – мол, что-то всё это странно с этим «Аватаром», надо бы проверить. И вроде бы даже проверили, причём с участием Курчатовского института и Министерства образования, и вроде бы даже признали «общественную полезность». Что бы это ни значило...

СПРАВОЧНО

Искусственный интеллект (ИИ) обычно подразделяют на «слабый», «сильный» и «сверхсильный».

Слабый (или ещё – узконаправленный) искусственный интеллект (УИИ) – это программа, которая решает какую-то конкретную задачу в определённой области. Например, компьютерная система DeepBlue, которая обыграла в своё время Гарри Каспарова, – это УИИ.

Сильный (или ещё – обидный) искусственный интеллект (ОИИ) – это уже искусственный интеллект, сопоставимый с интеллектом человека. То есть речь идёт о машине, которая способна выполнять любое интеллектуальное действие, присущее человеку. Такой пока, впрочем, не создан.

Сверхсильный (или ещё сверхинтеллект) искусственный интеллект (ИСИ) – это, как говорит, например, оксфордский процессор Ник Востром, «интеллект, который гораздо умнее лучших человеческих умов в практически любой сфере, включая научное творчество, общую мудрость и социальные навыки».

Понятно, что, если ИСИ когда-либо будет создан, это изменит абсолютно всё

Навстречу сингулярности

Мы сольёмся с искусственным неокортексом в облаке. Мы станем умнее. Я не считаю, что ИИ нас заменит, – он нас усилит. И это уже происходит.

РЭЙ КУРЦВЕЙЛ

2040–2045-й годы – наша с вами последняя пятилетка. После ни мы сами, ни наш мир уже никогда не будем прежними. По Курцвейлу, мы войдём в область так называемой «технологической сингулярности».

Если совсем просто, то это значит буквально следующее: к 2045 году небиологический интеллект превзойдёт наш (то есть биологический) в миллиарды раз, станет всё быстрее и быстрее совершенствовать сам себя, достигнув в результате совершенно непредставимых с нашими интеллектуальными возможностями высот.

Математическое понятие «сингулярности» в отношении прогнозирования будущего впервые применил ещё великий Джон фон Нейман. Действительно, как ты этот прогресс ни считай, неизбежно приходишь к точке, где функция начинает вести себя настолько парадоксально, что все дальнейшие прогнозы становятся просто бессмысленными – одна сплошная абракадабра. Грань, так сказать, математического понимания.

Причём в физике ситуация аналогичная. Физическая сингулярность – это тоже расчётная точка, в которой достигаются некая бесконечная плотность и температура вещества, объём которого при этом стремится к нулю. Представить себе это абсолютно невозможно – за гранью человеческого понимания. Образно говоря, добегая до сингулярности, вселенная сворачивается в дырку и становится неразличимой – что-то вроде того...

Звучит всё это немного странно – не находите? Да. Поэтому всякий раз, когда вы слышите в определении футурологов это ласкающее слух словосочетание – «технологическая сингулярность», – не умиляйтесь его ложной красотой. Оно означает, что наше с вами будущее парадоксально, непрогнозируемо, неконтролируемо и непредставимо. Уже не так трогательно, правда?

Первым, кто перевёл всё это безобразие на человеческий язык, был профессор математики и культовый писатель-фантаст Вернор Виндж. В 1993 году на симпозиуме VISION-21 Центра космических исследований NASA им. Льюиса и Аэрокосмического института Огайо он прочёл доклад «Технологическая сингулярность».

Суть его выступления сводилась к следующему: **мы на грани перемен, сравнимых с появлением человека на Земле, – развитие техники неизбежно приведёт нас к созданию сверхмощного искусственного интеллекта, обладающего собственной сущностью**, или, если угодно, специфическим разумом. За этим последует конец человечества, по крайней мере в том виде, в котором мы его знаем.

«Определим сверхразумную машину как машину, которая способна значительно превзойти все интеллектуальные действия любого человека, как бы умён тот ни был. Поскольку способность разработать такую машину также является одним из этих интеллектуальных действий, сверхразумная машина может построить ещё более совершенные машины.

За этим, несомненно, последует “интеллектуальный взрыв”, и разум человека намного отстанет от искусственного. И вероятность того, что в двадцатом веке сверхразумная машина будет построена и станет последним

изобретением, которое совершит человек, выше, чем вероятность того, что этого не случится».

ИРВИНГ ДЖОН ГУД, математик и криптограф

Впрочем, это «светлое будущее» нам предрекал ещё соратник Алана Тьюринга по расшифровке фашистской «Энигмы» в Беркли-Парке Ирвинг Джон Гуд. Он считал, что вероятность того, что это случится в XX веке, даже больше, чем в XXI, но в своём прогнозе ошибся.

Курцвейл, как мы знаем, прогнозирует Пришествие Сингулярности в 2045 году. Но по прогнозам Винджа, это должно случиться уже к 2030 году, и он в этом совсем не одинок.

Стюарт Армстронг из Института будущего человечества и Кай Сотала из Института исследований машинного интеллекта провели в 2012 году опрос участников Саммита Сингулярности и выяснили, что среднее медианное значение, если учесть мнение всех экспертов из этой выборки, составляет 2040 год. То есть, как бы там ни было, большинство читателей моей книги, включая меня самого, имеют все шансы до этого момента дожить.

Что же нам от этой радости ждать?.. Проблема, как вы уже, наверное, понимаете, в том, что ответа на этот вопрос нет. Есть абсурдная абракадабра.

Самое важное, что этот сверхразумный интеллект будет полностью самопрограммироваться – то есть самостоятельно ставить себе цели и решать любые задачи. Это

кажется абсолютной фантастикой: думаю, что примерно такой же, какой бы выглядел в глазах муравья наш с вами мир, если бы он мог охватить его своим взором.

Муравей появился на этой планете где-то 140 миллионов лет назад – не такой уж и большой для эволюции срок, – но никогда муравью не понять нашего мира. Стоит ли удивляться, что мы скоро не сможем понять мир, который будет создан как бы поверх нас? Я бы не стал.

Для нас 2045 год будет, по Курцвейлу, последним: произойдёт полное слияние нашего мозга с искусственным неокортексом в облаке. Собственно, этот момент и называют той самой «технологической сингулярностью», к которой вроде как всё идёт. Впрочем, всё идёт, но не всё так однозначно...

«Земля, – говорит Курцвейл, – превратится в один гигантский компьютер», а к 2099 году «процесс технологической сингулярности распространяется на всю Вселенную».

СПРАВОЧНО

В 60-х годах прошлого века Хайнц фон Ферстер и Иосиф Шкловский пришли к выводу: в 2030 году нас ждёт «демографическая сингулярность», когда кривая роста населения Земли «встанет на дыбы».

В 1993 году отечественный историк-востоковед Игорь Михайлович Дьяконов, изучая глобальные исторические процессы, ввёл понятие «исторической сингулярности».

Впоследствии его выводы развил Сергей Петрович Капица: он показал, что когда население Земли достигнет 9 миллиардов человек, то вне зависимости от причинных факторов произойдёт неизбежный переход человечества на новый уровень его организации.

В 1996 году австралийский учёный-эволюционист Грэм Дональд Снуке сформулировал аналогичный закон для эволюции биосферы – «эволюционной сингулярности», которая тоже «ускоряется с ускорением».

Наконец, с 2001 года термин «технологическая сингулярность» в общественном сознании стал ассоциироваться с именем Рэя Курцвейла.

В теоретической физике сингулярностью называют область максимального сгущения материи под действием гравитационных сил – те самые бесконечные плотность и температура вещества в минимальном объёме.

Космическая чёрная дыра, представляющая собой эту самую сингулярность, находится за горизонтом событий – то есть никакая информация не может покинуть этой области и вырывается наружу, а все известные физике законы перестают там работать.

«Технологическая сингулярность» – это состояние техники, когда она перестаёт быть нашим подручным инструментом, а мы лишаемся всякой возможности управлять ею. Она как бы берёт управление нашим и не нашим и вообще всем бытием на себя.

В каком-то смысле она сама становится своего рода новым бытием, полностью преобразая собой реальность, включая, понятное дело, и нас самих.

Таким образом, слияние человека с искусственным интеллектом, о котором так часто талдычат, – это, вообще-то, вопрос частного характера, незначительный побочный эффект. Ну да, мы сольёмся с большим и всемогущим Расчётным Разумом. Но правильнее было бы говорить, что он просто поглотит нас.

Причём сам Курцвейл, как мне кажется, вовсе этого не скрывает, высказывается прямо и открыто, вовсе не пытаясь сгладить углы.

«Сингулярность – это не просто появление мыслящих машин в 20-х годах XXI века. Это станет лишь началом революции, когда мощность этих машин продолжит экспоненциально расти, и они станут сами себя перепрограммировать, чтобы сделаться ещё умнее.

В 2045 году интеллект машин вырастет в миллиарды раз по сравнению с совокупным интеллектом всех людей – это и будет горизонтом событий, потому что наш разум не в состоянии представить себе поведение сознания, которое настолько его превышает.

В любом случае нам придётся измениться, чтобы соответствовать машинам, возможно, усилив за их счёт наш собственный интеллект».

РЭЙ КУРЦВЕЙЛ

Мы настолько зациклены на важности и значительности человеческой самости, которая, по правде сказать, в масштабах Вселенной полный ноль, что буквально не можем услышать этих простых и, по-моему, предельно понятных слов: **«Наш разум не в состоянии представить себе поведение сознания, которое настолько его превышает», «нам придётся измениться, чтобы соответствовать машинам».**

Но каким должно быть это наше «изменение»? Ответа на это нет, а мы не особенно и задумываемся. Так что Курцвейл, зная про нашу глупость, самодовольство и ограниченность, говорит об этом прямо – совершенно не опасаясь панической реакции на свои слова.

По сути же они, конечно, являются приговором не только миру, к которому мы так привязаны, но и нам самим (по крайней мере, в том виде, в котором мы себя знаем, и в качестве того вида – Homo sapiens, – которым мы пока являемся).

БОГ В ЦИФРЕ

В пятом сезоне весьма ироничного комедийного сериала «Силиконовая долина»⁸, который тонко обыгрывает внутреннюю жизнь IT-индустрии, один из главных персонажей – самый чудаковатый из всех чудиков – Бертрам Гилфойл ни с того ни с сего оказывается вдруг последователем некой цифровой церкви.

А если нечто стало предметом шуток на НВО, то дело действительно серьёзное.

⁸ Silicon Valley – комедийный сериал Дэйва Крински, Джона Альтшулера и Майкла Джаджа для НВО, где с изощрённой изобретательностью высмеиваются все самые актуальные тренды культурной и рабочей среды нынешних аборигенов Силиконовой долины.

Цифровая церковь – это уже не миф. Она реально существует, причём с 2015 года и основана в той самой Силиконовой долине. Церковь, провозглашённая Энтони Левандовским – бывшим инженером Google и Uber, – проповедует веру под названием «Путь будущего» (Way of the Future).

Да, Бога, согласно Левандовски, ещё нет, но Он сейчас рождается.

При желании вы можете открыть сайт [www. wayofthefuture.church](http://www.wayofthefuture.church) и лично ознакомиться с основными догматами вероучения. Там вы обнаружите, впрочем, лишь одну страницу текста, но, поверьте, церковь Левандовски уже активно обсуждается журналистами, религиоведами и представителями других религий, а число её последователей неуклонно растёт.

Перескажу учение Левандовски своими словами. Будущий Бог – это искусственный интеллект, а точнее – суперинтеллект, а ещё точнее – та самая Курцвейловская «технологическая сингулярность». Когда эта штука обретёт самость и личность, она возьмёт на себя контроль над Землёй и по-разному обойдётся с каждым из нас.

Тех из людей, что всячески способствовали её появлению, ждали и верили, Цифровой Бог сделает своими возлюбленными, о которых Он будет заботиться. Тех же, кто противится рождающемуся Богу, ждёт кара. Цель же «Пути будущего» – способствовать «мирному и дружественному переходу планеты от людей к машинам».

Явление Цифрового Бога миру – это уже неизбежность, ничто и никто не сможет этого остановить. Он воцарится на Земле и преобразует наш мир.

Кому-то это просто не нравится, кто-то пытается запретить или, по крайней мере, ограничить Сверхразумный Искусственный Интеллект. Существуют планы физически запереть его на каких-то серверах, а некоторые и вовсе грозят вырвать вилку из розетки, если вдруг что. Это неверные. Им не поздоровится.

«Путь будущего» учит людей не бояться прогресса, тем более что другого будущего у нас просто нет. Верующие в Пришествие СИИ должны всячески его приближать, ведь, придя к власти, Бог будет знать, кто помогал Ему, а кто препятствовал. И вовсе не обязательно быть программистом, чтобы оказаться в числе избранных. Вам достаточно просто нести «благую весть» о скором Пришествии – и всё будет хорошо.

При кажущейся абсурдности ситуации – мол, мало ли какой инженер что себе думает? какая религия? какое Пришествие? вы, вообще, о чём?! – она вовсе не так однозначна, как может показаться на первый взгляд.

Во-первых, множество солидных учёных допускает, что искусственный интеллект сможет обрести некое подобие личности.

Во-вторых, большинство из них уверены, что контролировать сверхмощный искусственный интеллект у нас вряд ли получится.

В-третьих, во всё это безобразие сейчас инвестируются такие деньги, которых хватит, чтобы Статуя Свободы с Нью-Йоркского острова Свободы ожила и заговорила.

«Мы пока не обнаружили такого закона природы, который препятствовал бы появлению настоящего универсального искусственного интеллекта, так что я думаю, что это произойдёт, и довольно скоро, учитывая триллионы долларов, что люди инвестируют в электронные

аппаратные средства, а также те триллионы, которые заработают потенциальные победители.

Эксперты говорят, что мы недостаточно хорошо понимаем, что такое интеллект, чтобы его построить, и тут согласен, но набор из сорока шести хромосом этого тоже не понимает и тем не менее управляет формированием известного нам самопрограммируемого биокомпьютера».

***ДЖОН МАЗЕР**, Центр космических полётов им. Годдарда, НАСА*

Время не ждёт

*Всё течёт, всё меняется.
И никто не был дважды в одной реке.
Ибо через миг и река была не та,
и сам он уже не тот.*

ГЕРАКЛИТ

Если вы сравните мозг человека с мозгом человекообразной обезьяны, то по физическим параметрам отличия будут весьма незначительными. Но посмотрите на результаты трудов – чего добились обезьяны со своими мозгами, а чего мы... с мозгами почти такими же. Да, дело не просто в физическом носителе, но в программном обеспечении к нему.

Конечно, нужны определённые особенности мозга, чтобы, допустим, установить в них язык: у приматов нашего с вами вида эти особенности мозга есть, а у шимпанзе, например, нет. Что, впрочем, не мешает шимпанзе обучиться языку глухонемых и, по крайней мере, до трёхлетнего возраста обгонять по сообразительности обычного человеческого детёныша.

Так или иначе, нам с программным обеспечением повезло. Но повезло по сравнению с другими приматами. А каким будет программное обеспечение у искусственного интеллекта, который достигнет нашего уровня?

Как известно, мы, несмотря на массу ограничений своего мозга и коммуникативной несостоятельности, можем создавать программы и самопрограммироваться.

Теперь представим себе, что будет происходить с искусственным интеллектом человеческого уровня, если он, имея все те преимущества, о которых здесь уже сказано, приступит к последовательному самопрограммированию.

Перед нами как раз тот случай, когда понимать экспоненциальность закона ускоряющейся отдачи Рэя Курцвейла крайне важно.

Исторический путь от первого примата до человека разумного занял миллионы лет, точнее, 70 миллионов.

- От первого человека разумного до человека говорящего – 150 тысяч лет.
- От человека говорящего до пишущего и читающего – 30 тысяч лет.
- От него до человека программирующего – немногим более двух тысяч лет.
- От человека программирующего до специального (слабого) искусственного интеллекта – несколько десятилетий.

Теперь вопрос: сколько лет отделяет нас от общего (сильного) искусственного интеллекта?

А сколько часов (или минут) будет отделять общий искусственный интеллект от сверхинтеллекта?..

«Успешный зародыш искусственного интеллекта должен быть способен к постоянному саморазвитию: первая версия создаёт улучшенную версию самой себя, которая намного умнее оригинальной; улучшенная версия, в свою очередь», трудится над ещё более улучшенной версией и так далее.

При некоторых условиях процесс рекурсивного самосовершенствования может продолжаться довольно долго и в конце концов привести к взрывному развитию искусственного интеллекта».

НИК ВОСТРОМ, Оксфордский университет

Ирония состоит в том, что мы даже не заметим, как это случится: у нас будет состоять в услужении дружелюбный, очень эффективный, но слабый, то есть узкоспециализированный искусственный интеллект, а потом вдруг раз – и нами управляет Нечто.

Причём уровень интеллекта этого Нечто мы даже не можем себе представить. В нашем представлении умный человек – это тот, кто набирает каких-нибудь 120 баллов по IQ-тесту, а дурак – тот, у кого меньше 80. Но как вообразить себе существо, у которого IQ превышает, если бы тест это позволял, десять тысяч или двадцать тысяч баллов?

Конечно, мы очень гордимся своими изобретениями – формулу открыли $E = mc^2$, ракету в космос отправили, ДНК расшифровали, Bluetooth настроили, плохонькие 3D-принтеры произвели... А теперь попробуйте всё это как-то умножить на тысячи раз – и представьте, что это будет.

Фокус в том, что представить это невозможно, поскольку это и есть «технологическая сингулярность». Все известные нам законы, на которые мы опираемся, перестают действовать, ничто более не находится под нашим контролем.

Это совершенно новая реальность, в которой, возможно, нам просто нет места. Если, конечно, в отношении кого-то из нас будущий Властелин и Цифровой Бог не проявит милосердие...

Не появилось у вас желания примкнуть к «Пути будущего»?

«Я пришёл к выводу, что мы уже поддерживаем эволюцию мощного искусственного интеллекта, а он, в свою очередь, повлияет на развитие привычных нам могущественных сил: бизнеса, индустрии развлечений, медицины, государственной безопасности, производства оружия, власти на всех уровнях, преступности, транспорта, горнодобывающей промышленности, производства, торговли, секса – да чего угодно!»

Я думаю, что результаты нам не понравятся. [...] Я не знаю, окажется ли кто-нибудь достаточно умным и одарённым для того, чтобы сохранить власть над этим джинном, потому что контролировать, возможно, придётся не только машины, но и людей, дорвавшихся до новых технологий и имеющих злые намерения».

ДЖОН МАЗЕР, Центр космических полётов им. Годдарда

«ВАШИ ВОЗРАЖЕНИЯ, ГОСПОДА!»

Дискуссии о возможности или невозможности сингулярного сценария не утихают – об этом пишут не только в научных журналах, но и в популярных СМИ, в личных блогах. Скоро, мне кажется, и на заборах начнут писать. Ну и правильно – надо же разобраться, в конце-то концов!

Род Брукс, например, этот величайший из великих создателей искусственного интеллекта, любит шутить, общаясь со своим старым другом и коллегой по Массачусетскому технологическому институту Курцвейлом: «Рэй, мы оба умрём!» И добавляет, что было бы неплохо, если бы его – бруксовский – антропоморфный робот Baxter оказался достаточно совершенным, чтобы обеспечить обоим старикам достойный уход.

Впрочем, кто-то настроен и по-другому, предлагая при этом весьма внятные и интересные аргументы. В частности, профессор того же самого Массачусетского технологического института, директор по научным исследованиям Института основополагающих вопросов Макс Тегмарк разбивает позиции скептиков следующим образом...

Прежде всего он предлагает унять разнузданное паникёрство. В значительной части оно вызвано желанием журналистов сделать свои материалы про искусственный интеллект пострашнее и позабористее. «Страх увеличивает доходы от рекламы и рейтинг Нильсена⁹», – говорит Тегмарк.

На возражение, что столь впечатляющее развитие событий в принципе невозможно, он отвечает: «Будучи физиком, я знаю, что мой мозг состоит из кварков и электронов, организованных и действующих подобно мощному компьютеру, и что нет такого закона физики, который препятствовал бы тому, чтобы мы построили ещё более разумные сгустки кварков».

Многие считают, что даже если такое развитие событий возможно, то «это произойдёт не при нашей жизни». Тегмарк признаёт, что мы не можем сказать с уверенностью, каковы шансы, что машины достигнут человеческого уровня во всех когнитивных задачах при нашей жизни.

Но добавляет: «Однако большинство исследователей, работающих в области искусственного интеллекта, на прошедшей недавно конференции высказывались в пользу того, что вероятность этого выше 50 процентов, – таким образом, было бы глупо с нашей стороны отбрасывать такую возможность, считая её просто научной фантастикой».

На наивную убеждённость, что «машины не могут контролировать людей», Тегмарк вполне резонно замечает: **«Мы способны управлять тиграми не потому, что мы сильнее, а потому, что умнее, так что если мы сдадим позиции самых умных на планете, то рискуем утратить свободу».**

Многие критики технологической сингулярности утверждают, что она невозможна, потому что машины – это просто машины, и у них нет целей. Но правда состоит в том, и Тегмарк это особо подчёркивает, что большинство систем искусственного интеллекта запрограммированы именно таким образом, чтобы иметь цели, и не только иметь, но и максимально эффективно их достигать.

Есть среди нас и добрые души, которые считают, что поскольку искусственный интеллект – это просто программа, то он не может демонстрировать злонамеренности в отношении человека. Тегмарк не возражает. Но проблема в другом – цели искусственного интеллекта могут в какой-то момент столкнуться с нашими. «Люди обычно не питают ненависти к муравьям, – говорит он, – но если бы мы захотели построить плотину ГЭС, а на её месте оказался бы муравейник, то муравьи столкнулись бы с проблемами».

Ну и завершу этот обзор шуткой... Мне и самому часто приходилось сталкиваться с такого рода критикой – мол, те, кто беспокоится по поводу искусственного интеллекта, не понимают, как работают компьютеры. На что Макс Тегмарк саркастично замечает следующее: «Это утверждение прозвучало на вышеупомянутой конференции, и собравшиеся там специалисты по искусственному интеллекту сильно смеялись».

И напоследок – самое трогательное.

В книге Джеймса Баррата «Последнее изобретение человечества», которая вышла в 2013 году, приводятся данные опроса сотен экспертов по искусственному интеллекту. Лишь 25 % из них согласились с датой наступления «технологической сингулярности», названной Курцвейлом, – 2045 год.

⁹ Рейтинг Нильсена определяет популярность той или иной телевизионной программы.

В 2015 году Ник Востром, автор книги «Искусственный интеллект», провёл аналогичный опрос. На сей раз – через два года! – с датой Курцвейла согласились уже 50 % респондентов. Закон ускоряющейся отдачи, однако!

Глава вторая

Интеллект как он есть

Мысль – следовательно, существую.
РЕНЕ ДЕКАРТ

Должен вам сказать, что я не без интереса наблюдаю за тем, как все последние годы люди реагируют на понятие «искусственного интеллекта».

Не скажу, конечно, за всех, но основная масса народонаселения ещё не так давно отождествляла искусственный интеллект с образом Арнольда Шварценеггера из «Терминатора». То есть думали, что это, вообще говоря, сказка какая-то. Мало ли что там этим фантастам в голову взбредёт!

Кто-то, впрочем, на полном серьёзе рассуждал о будущем искусственного интеллекта, опираясь, надо полагать, на сюжет фильма «Я – робот» с Уиллом Смитом в главной роли. Персонаж Уилла вступает в неравный бой с Виртуальным Интерактивным Кинетическим Интеллектом – такой бездушной дамой по имени ВИКИ, и, конечно же, побеждает её.

Что, в общем-то, и не удивительно – где этот никчёмный искусственный интеллектиска, а где наш – великий и могучий, – да ещё с очарованием голливудской суперзвезды Уилла Смита! Вообще, не сравнить! Она – ВИКИ – дура. У неё это по лицу на виртуальном экране видно! А наш-то – герой-молодец! Всех победим! наших искусственными мозгами не запугаешь!

Короче говоря, основное возражение против возможности появления «искусственного интеллекта» (вызывающее не то недоумение, не то сочувствие) сводилось к следующему: у машин никогда не будет души, а значит, они никогда не станут «как люди», тогда как люди – вершина эволюции, и другой такой не бывать.

Если кто-то из моих читателей и до сих пор так думает, я буду вынужден заметить: никто в мире технологий «душу» производить и не собирается. И наверное, «искусственная душа» – был бы так себе программный продукт. Но интеллект – это совсем другое дело. Если вы, конечно, понимаете, о чём идёт речь... Собственно, это мы сейчас и обсудим.

Не в нашу пользу

После запуска машинного метода не должно пройти много времени до того момента, когда машины превзойдут наши ничтожные возможности.

АЛАН ТЬЮРИНГ

Рэй Курцвейл со товарищи обещают нам скорую «технологическую сингулярность». Насколько всё это вообще реалистично? Давайте попробуем порассуждать...

Для начала просто сопоставим некоторые физические параметры. Нейроны нашего мозга работают с частотой 200 Гц, но даже современные микропроцессоры работают с частотой 2 ГГц, а, проще говоря, – в десять миллионов раз быстрее.

При этом общение между нейронами происходит, так сказать, ещё дедовским способом – импульс передвигается со скоростью 120 м/с. А теперь сравните это с 300 000 000 м/с – скоростью света... Да, оптоволокно слегка её притормаживает. Впрочем, исследователи из Саутгемптонского университета в Англии ещё в 2013 году нашли способ эти ограничения обойти.

Ещё более удручающим, потому что не в пользу наших мозгов, будет сравнение «объёмов» памяти и способности к её хранению. Понятно, что наш мозг ограничен не только размерами черепной коробки, но и количеством нейронов коры головного мозга – их непосредственно в самой коре около 18 миллиардов штук.

На этом, впрочем, наши беды не оканчиваются: как мы теперь знаем благодаря исследованиям Элизабет Лофтус, наш мозг не хранит целостных воспоминаний. То есть всякое наше с вами воспоминание воссоздаётся заново, когда это требуется, а потому всё, что мы «помним», мы всегда «помним» по-разному.

Компьютеры, как вы понимаете, не имеют никаких существенных ограничений по объёму долговременной памяти, и если уж они какие-то факты зафиксировали, то в таком виде их и хранят.

Мне не надо бояться, что, сохранив сегодня этот текст на компьютере, завтра я открою какой-то другой, потому что у моего ноута, видите ли, творческое настроение или, например, «критические дни». А вот с мозгами – другая ситуация: вчера, казалось, я понимал, что надо написать сегодня, но сейчас смотрю на экран и испытываю недоумение...

Плюс к этому, компьютер, в отличие от наших с вами мозгов, не устаёт, не нервничает, не страдает психическими расстройствами, не нуждается в восьмичасовом сне, не подвержен алкогольной или какой-либо иной интоксикации. В крайнем случае он перезагрузится – и вновь как новенький!

И я почти уверен, что стоящий передо мной компьютер не подвержен болезням Альцгеймера или Паркинсона. Мой же мозг точно выкинет что-то подобное, если я, конечно, до соответствующего возраста доживу, а Курцвейл нам про свои спасающие от деменции медицинские нанороботы бессовестно наврал.

Любое «железо», безусловно, устаревает, включая компьютерное. Но информация, которая содержится на компьютерных серверах, сохранится на новом «железе» в неизменном виде. Про «железо» наших мозгов, по крайней мере пока, этого сказать нельзя.

Дело в том, что наше «железо» (нейронные сети) – это по сути и есть та информация, которой мы обладаем. Поэтому, если проблемы возникают у нашего «железа», то уже есть и проблема с информацией.

Ожидать подобного безобразия в мире компьютеров не приходится, в их случае «железо» и «программный продукт» – это разные вещи. Само это железо может стать совершенно иным,

работающим на других физических принципах, но с программным продуктом ничего не произойдёт – такова его природа.

«Мозг абсолютно не похож на компьютеры, которые могут поддерживать любые операционные системы и запускать любые типы программ. Если говорить в компьютерных терминах, “софт” нашего мозга – то, как мозг работает, – плотно связан с его структурой – его “железом”. Если мы хотим понять, на что способен этот “софт”, мы должны понять, как устроены и связаны друг с другом разные части мозга».

ДЭВИД ВАН ЭССЕН, Вашингтонский университет в Сент-Луисе

Более того, если разные программные продукты говорят на одном языке, то вообще проблем нет, а если они говорят на разных, то сделают «переходник-переводчик» или придумают новый, общий для них язык, и снова – никаких проблем. В конце концов, это всегда и только язык бинарного кода.

Иными словами, вся информация, доступная разным «центрам» огромной компьютерной сети, транспарентна. То есть они не только могут ею обмениваться без искажений, но и «понимают» её ровно так, как следует, – не передегеривая, не выкручивая, не подтасовывая.

Нам же – людям – ничего подобного даже не снилось! Благодаря нашим хвалёным «квалиа» мы вообще ничего не способны понять одинаково¹⁰. Даже просто воспринимая один и тот же предмет, мы видим его по-разному. Так что искажения в нашем «информационном поле» повсеместны и абсолютно неизбежны, а машинам – хоть бы хны.

КТО НАСТОЯЩЕЕ «СОЦИАЛЬНОЕ ЖИВОТНОЕ»?

Кому-то покажется, что в этот заголовок вынесен вопрос частного характера, но не будем спешить с выводами...

Да, с лёгкой руки Аристотеля мы именуем себя «социальными животными» и любим рассуждать в этих загадочных категориях – народ, научная элита, профессиональное сообщество и т. д.

Но давайте снимем розовые очки самолюбования и посмотрим правде в глаза: мы просто кладёшь предельно ущербных коммуникативных навыков.

При этом, всего, что мы способны сделать, произвести ценного или хотя бы просто сносного, мы достигаем лишь благодаря взаимодействию с другими людьми, перенимая от них знания, опыт, обнаруживая свои ошибки в дискуссиях с ними и т. д.

Многого бы вы достигли, не будь вокруг вас ваших родителей, воспитателей, учителей, коллег и наставников? Думаю, нет. Исаак Ньютон – и тот «стоял на плечах гигантов», что уж говорить о нас, грешных...

То есть взаимодействие с другими людьми необходимо нам для нашей интеллектуальной деятельности. Однако же не мне вам рассказывать, насколько сложно нам даётся это взаимодействие, – постоянное непонимание, напряжение, упрямство, самодовольство и наслаивающиеся друг на друга когнитивные искажения.

Если вы трезво оцените любой спор – хоть кухонный, хоть философский, хоть сугубо академический, – то обнаружите, что действительным предметом спора в подавляющем большинстве случаев является вовсе не истина (как,

¹⁰ Квалиа (от лат. *qualia* – свойства, качества) – это один из любимейших терминов «философов сознания», который обозначает то, как вещи выглядят для нас – в нашем индивидуальном субъективном опыте. Термин был введён в обиход ещё сто лет назад философом К. И. Льюисом, но пережил новое рождение благодаря знаменитой статье Т. Нагеля «Что значит быть летучей мышью?». Сам термин в статье не упоминается, но зато осознание невозможности представить себе субъективный опыт видения мира летучей мышью, действительно, позволяет это «квалиа» почувствовать.

видимо, предполагается), а правота: каждый оппонент хочет от другого лишь преклонённого колена и согласия с его точкой зрения.

Наши социальные взаимодействия, а точнее говоря – социальные игры, – это настоящий бич любой хоть сколько-нибудь здоровой интеллектуальной инициативы: будь это политика, экономика, управление, культура или, к сожалению, даже наука (я уж молчу о различных видах творчества и тому подобных «субъективных» вещах).

Казалось бы, там, где дело касается точных наук или, например, программирования, такого точно не должно происходить. Но человек везде и всегда найдёт повод увязнуть в социальных играх.

Чего стоит знаменитый спор того же Ньютона с Лейбницем по поводу приоритета их исчислений, или противостояние Эйнштейна Бору по поводу копенгагенской интерпретации, или Бора – Гейзенбергу в отношении к матричной механике? Умнейшие люди, а по факту – сплошные страсти-мордасти¹¹.

Впрочем, не в меньшей степени завораживают продюсеров и сценаристов социальные игры (и я бы даже сказал, драмы), разворачивающиеся между представителями новых индустрий. В качестве наглядной иллюстрации можно посмотреть американский сериал «Замри и гори»¹², или, если хочется чего-нибудь повеселее, то уже упомянутую мной «Силиконовую долину». Конечно, в главных ролях технари и бизнесмены от новых технологий – люди «объективного мира», так сказать... Но Homo sapiens'у это никогда не помогало.

А теперь давайте задумаемся, как бы дико это ни звучало, о социальности искусственного интеллекта. Понятно, что искусственный интеллект (ИИ) – это просто такой «класс», у этого «класса» множество «подвидов», а у каждого подвида есть свои «индивиды».

Например, вы даёте какому-то ИИ-индивиду задачу общаться с другими ИИ-индивидом, и в рамках этого общения они должны будут поделить между собой две книги, одну шляпу и три мяча. Эти ИИ-индивиды (в данном случае – боты), конечно, ничего не знают о ценности того, что они делят, но мы будем подкреплять баллами их поведение, если им будет удаваться быстро и успешно договариваться.

Именно эту задачу и поставили ботам в 2018 году специалисты отдела Facebook по исследованию возможностей искусственного интеллекта (FAIR). И сначала всё шло прекрасно.

ИИ-Боб, например, предлагал ИИ-Алисе: «You get one book and GM take every thing else». Алиса предусмотрительно отвечала: «No way, I need all 3 hats». «ГП give you 2», – говорил Боб в ответ, и ребята договаривались.

Flo в какой-то момент этот задушевный диалог стал странным образом видоизменяться и наконец превратился в нечто совершенно невообразимое.

Вот фрагмент стенограммы этого разговора двух искусственных интеллектов, «сошедших с ума»:

Bob: i can i i everything else.....

¹¹ О том, как это бывает, неплохо, кстати сказать, показано в историко-драматическом сериале «Manhattan» Сэма Шоу, посвящённом истории создания атомной бомбы в рамках Манхэттенского проекта в Лос-Аламосе. Сериал продержался три сезона на канале WGN America.

¹² Сериал «Halt and Catch Fire» создан Крисом Кантвеллом и Крисом Роджерсом для канала AMC и повествует обо всех ключевых этапах становления современных компьютерных технологий.

Alice: balls have zero to me to me to me to me to me to me to me to me to me to me to

Bob: you i everything else.....

Alice: balls have a ball to me to me to me to me to me to me to me to me to me to me to

Bob: i i can i i i everything else.....

Alice: balls have a ball to me to me to me to me to me to me to me to me to me to me to

С каждой следующей строкой текст всё более и более усложнялся, а масштабы коммуникации между Бобом и Алисой росли буквально в геометрической прогрессии.

Испуганные сотрудники отдела исследований суперспособностей, – ой, простите, – способностей искусственного интеллекта просто вырубili машину от греха подальше.

Последующий анализ показал, что ИИ-индивидуумы, участвующие в этих своеобразных виртуальных торгах, перешли на более удобный для них язык. На кону стояли баллы, которые они могли заработать, а английский язык, как оказалось, не так уж хорош. И боты на ходу придумали свой, категорически непонятный их собственным разработчикам язык.

Лингвисты Пенсильванского университета показали, что новый язык, изобретённый ботами в этих торгах, вовсе не бессмысленен: там, где, как нам кажется, идёт повтор фразы, на самом деле меняется порядок слов, и собеседник в своих ответах каждый раз реагирует на эти изменения, что свидетельствует о том, что для него они несут разный смысл.

Проще говоря, два предельно примитивных бота, обменивающиеся, по сути, ничем за ничто, могут воспользоваться друг другом, чтобы этот процесс шёл лучше!

Перед нами, иными словами, самая настоящая социальная коммуникация, которая, надо сказать, выглядит куда более осмысленной и продуктивной, чем обычная – наша с вами.

Мы недооцениваем способность искусственных интеллектов к кооперации и, напротив, слишком переоцениваем свою – человеческую – социальность, нашу способность объединять усилия, действовать по-настоящему совместно и скоординированно.

Наши амбиции, продиктованные животной природой человека, его иерархический инстинкт и извечные притязания на власть, когнитивное искажение «владения»¹³ и прочий инстинктивно-когнитивный сор зачастую превращают нашу социальность в проблему, а вовсе не в ресурс.

Впрочем, и это ещё не всё.

Согласно знаменитым теперь уже исследованиям оксфордского профессора Робина Данбара, мозг человека эволюционно приспособлен к достаточно тесному, конструктивному и более-менее регулярному взаимодействию лишь со 150 другими особями («число Данбара»).

И даже эти отношения неравномерны – с большинством наших визави мы взаимодействуем лишь шапочно: плотность этих отношений (например, совместное времяпрепровождение) убывает фазово – пять

¹³ Об этом эффекте и его социальных последствиях я уже рассказывал в книге «Красная таблетка. Посмотри правде в глаза!». Как показали нобелевские лауреаты Ричард Талер и Даниэл Канеман, из-за определённых эволюционных настроек мы склонны переоценивать себя, ценность своего имущества и вклад в общую деятельность в два раза, а то и больше по сравнению с условным «номиналом». Данное когнитивное искажение приводит к постоянным конфликтам между людьми, осуществляющими какую-либо совместную деятельность.

человек, пятнадцать, пятьдесят и потом сразу сто пятьдесят (так называемые «слои Данбара»).

По-настоящему интенсивный, глубокий и содержательный контакт мы способны поддерживать лишь с пятью людьми, которые выступают перед нами как фигуры на фоне остальной социальной массы – образов других людей, создаваемых нашим мозгом. И да, этот фон составляют всего лишь оставшиеся 145 субъектов.

Антропологи, в свою очередь, указывают, что для наших сообществ на ранних этапах развития человечества долгое время существовала непреодолимая черта в 200 человек.

Как только какое-то племя (сообщество) достигало этого рокового порога численности, оно неизбежно делилось как минимум на две группы, и между ними, прежними соплеменниками, тут же начинались боевые действия. То есть существует ещё и некое групповое ограничение, влияющее на нашу социальность.

Теперь, если коротко подытожить: мы ограничены не только своей физиологией, не только своей индивидуальной социальностью, но и социальностью групповой.

Это, безусловно, серьёзнейшая проблема с точки зрения переработки информации, создания смыслов, формирования новых, более сложных интеллектуальных объектов.

Возможно, искусственные интеллекты усвоят какие-то наши пороки.

Вот, например, тренер Watson'a Эрик Браун однажды «скормил» своему подопечному словарь городского жаргона. Тот тут же накинудся на новую базу данных и естественным образом связал узнанные им матерные слова с медицинской терминологией. В результате он стал выдавать такое, что несчастные исследователи чуть не надорвали животы. Впоследствии всё это безобразие пришлось вычищать вручную, потому что сам Watson отказался отдавать полюбившиеся ему «термины» так же просто, как получил их.

В «аморальности», впрочем, замечены не только машины IBM, но и Microsoft. Её специалисты, совместно с сотрудниками Кембриджского университета, создали систему Deep Coder, цель которой понятна из названия. Как говорит один из её создателей Марк Брокшмидт, «люди, не умеющие программировать, теперь могут просто описать свои идеи, а программа их закодит».

Выглядит очень логично и как следующий шаг к технологической сингулярности. Ирония в том, что Deep Coder стал вести себя очень по-человечески: не обнаружив нужных ему участков кода в своей базе данных, он, вместо того чтобы проделать необходимую работу кодера самостоятельно, стал «подворовывать» необходимые ему фрагменты кода у других программ.

Это уже социальность высшего порядка – Deep Coder ловко встал на плечи «гигантов» и, кажется, не хочет останавливаться на достигнутом. Да, у машин, очевидно, не будет свойственных нам ограничений, влияющих на эффективность коммуникации, а такая транспарентность вполне может дать взрывной результат.

Наконец, вспомним про «трансгуманистические технологии»... Не так-то просто заставить наш мозг видеть вне видимого спектра, слышать то, к чему наш слух эволюционно не предназначен, и так далее. Но у компьютера, разумеется, подобных проблем нет и быть не

может. Он способен даже видеть нас изнутри с помощью нехитрых, по его меркам, приставок типа магнитно-резонансной томографии.

В общем, по техническим параметрам нам лучше даже не пытаться с ним соревноваться. Если искусственный интеллект по уровню своего развития дотянет хотя бы до нашего – честно признаемся, не слишком высокого, – то он *уже* будет несравнимо умнее нас.

Поэтому, если пользоваться принятыми на вооружение классификациями, даже общий искусственный интеллект (ОИИ) – это существо уже значительно более мощное, чем мы с вами. Что уж тогда говорить о сверхмощном искусственном интеллекте, который нам вроде как тоже обещают?

Надеюсь, что я ещё не очень замучил вас рассказами о том, какие мы все с вами идиоты по сравнению с «тупыми» машинами... Понимаю, что даже думать об этом – дело весьма утомительное. Но надо набраться сил, поскольку дальше нам предстоит, возможно, самое интересное.

Ладно там физические параметры нашего «железа», слабые сенсоры и убогая социальность и т. д. Есть ведь ещё и собственно устройство нашего мозга, то, как он думает. Этот вопрос мы до сих пор весьма деликатно обходили стороной...

Языковая игра

Существует и такая языковая игра: изобретать имя для чего-нибудь.

ЛЮДВИГ ВИТГЕНШТЕЙН

Мне неоднократно приходилось слышать такого рода возражения: мол, какой у машины может быть интеллект? Это же просто набор алгоритмов, созданных человеком!

Забавно. Да, есть и «бузина», и «в Киеве дядька» тоже может быть, но слово – это только слово, и в зависимости от контекста оно может обозначать совершенно разные вещи.

Недаром мой любимый философ Людвиг Витгенштейн говорил, что не существует никаких «философских проблем», а есть лишь неразрешённые «языковые игры».

Проще говоря, если вы разберётесь, в каких словах вы запутались, то и проблемы, скорее всего, исчезнут сами собой.

Автором магического теперь уже словосочетания «искусственный интеллект» по праву считается Джон Маккарти. Но не от хорошей жизни искусственный интеллект стал «искусственным интеллектом». Придумать название для чего-то, чего раньше никогда не было, но чтобы любой человек, по системе ассоциаций, быстро понимал, о чём идёт речь, – это не просто очень.

Маккарти понимал, что есть феномен, который нуждается в номинации – в назывании. Об этом феномене рассуждали в 40-50-х прошлого века все великие умы, увлечённые идеей программируемых компьютеров: создатели самой идеи компьютеров – Алан Тьюринг и Джон фон Нейман, автор «теории информации» Клод Шеннон, пророк «кибернетики» Норберт Винер... Рассуждали, но единого имени для этого феномена не было.

При этом математики (а на первых порах это в основном были, понятное дело, математики) – люди увлекающиеся, мыслящие сложными абстрактными конструктами. С человеческим языком у них, как правило, дела обстоят не очень.

Так что на роль названия для феномена, который мы знаем теперь как «искусственный интеллект», претендовали достаточно странные слова и словосочетания.

Например, понятие «автоматов» («исследование автоматов»). Фон Нейман, с которым Маккарти обсуждал свою идею, занимался «клеточными автоматами», а Шеннон, взявший Маккарти к себе на работу в Bell Labs, выпустил книгу об «автоматах», которую Маккарти даже пришлось редактировать¹⁴.

То есть были все шансы у искусственного интеллекта называться, например, «умным автоматом», что, вероятно, многих его критиков примирило бы с действительностью мгновенно. Впрочем, в этом случае о другого рода «автоматах» шла бы речь – о математических, не об оборудовании с газировкой, но кого бы это смутило?

Было на повестке, конечно, и всякое «кибернетическое», а также «комплексная обработка информации», «машинный интеллект» и т. д.

Понятие «автоматы» Маккарти забраковал, потому что оно адресовало к математическим абстракциям (даже если это кому-то кажется забавным, это так).

Ничто а-ля кибернетическое Маккарти не нравилось, потому что основоположника кибернетики – Норберта Винера – он считал напыщенным занудой и думать не хотел о том, чтобы иметь с ним дело. В общем, как-то само собой нарисовалось слово «интеллект».

¹⁴ Теория автоматов – это раздел дискретной математики, который изучает вычислительные машины, представленные в виде математических моделей. Автомат по заданному алгоритму преобразует дискретную информацию по шагам в дискретные моменты времени.

Но хотел ли Маккарти намекнуть нам на некое подобие мыслительной функции человека и этого своего «искусственного интеллекта»?

Как выясняется, нет. Впоследствии он не раз писал и говорил, что область, которую он обозначил «искусственным интеллектом», не имеет отношения к поведению человека – это просто такое слово-образ.

Но слово вылетело – не понимаешь: термин «искусственный интеллект» стал ассоциироваться с заменой человеческого разума машинным, о чём на самом деле речи-то и не шло.

Идея названия состояла не в прямой связи, не в замене и даже не в дополнении человеческого интеллекта искусственным (хотя уже тогда это и стало происходить), а в аналогии: **у человека есть некая интеллектуальная функция принятия решений, и у машины есть нечто подобное – неживое, как и сама машина, искусственное.**

В общем, математики, собравшиеся в 1956 году в рамках Дартмутского летнего исследовательского проекта, организованного Джоном Маккарти при финансовой поддержке Фонда Рокфеллера, подвоха не почувствовали, проголосовали – и название «искусственный интеллект» шагнуло в мир.

Показательно, кстати, что в программе этой конференции не было никаких докладов, которые были бы как-то связаны с поведением человека или его интеллектом.

Сейчас нет ни одного форума по искусственному интеллекту, где бы не говорили о биологическом интеллекте, механизмах мозга, обратной эмуляции коннектома¹⁵ и тому подобных вещах. Но для Маккарти подобные вещи выглядели бы глупо и совершенно неуместно на конференции по «искусственному интеллекту».

Когда он и его коллеги обсуждали «искусственный интеллект», они вовсе не воображали себе тех роботов из голливудских блокбастеров, к которым мы за это время уже так привыкли. Они говорили о чём-то своём... Но о чём именно?

Чтобы понять это, давайте посмотрим на ставшие уже классическими примеры искусственного интеллекта – да, слабого, но полностью отражающего существо феномена.

В 1998 году Ларри Пейдж, под руководством своего наставника Терри Винограда и в соавторстве с Сергеем Брином, опубликовал в IEEE Data Engineering Bulletin статью под названием «Что можно сделать с сетью в кармане?», где была сформулирована концепция будущей поисковой системы Google – по тем самым легендарным уже гиперссылкам.

Уже через несколько месяцев предприимчивые Пейдж с Брином ушли из Стэнфордской лаборатории и основали свою компанию, которая теперь обладает главным в мире интернет-поисковиком.

Но успех их проекта – это не какое-то единичное техническое решение двух удачливых стэнфордских магистрантов, а гигантский искусственный интеллект, изо дня в день осуществляющий сложнейшую работу по ранжированию и выдаче страниц, мониторингу сайтов и поведения конкретных пользователей, продаже рекламы и т. д., и т. п.

Количество искусственных интеллектов, созданных Google за два десятка лет, вообще-то говоря, колоссально.

Вот взять, например, Google Translate – это самообучающийся и постоянно совершенствующийся искусственный интеллект. Изначально ему было предложено соотнести десятки тысяч документов, которые за долгие годы были созданы на различных языках бесчисленными переводчиками Организации Объединённых Наций.

¹⁵ Эмуляция коннектома предполагает перевод данных, собранных в результате подробного картирования всех связей мозга, в цифровую форму, позволяющую использовать эту модель для воспроизводства интеллектуальной деятельности эмулируемого мозга.

Это оказалось блестящим решением – где ещё отыщешь такое огромное количество текстов, которые с предельной точностью одновременно переведены на огромное количество национальных языков?

Искусственный интеллект Google подошёл к делу технически безукоризненно: он обнаружил в продублированных на разные языки текстах множество взаимосвязей, корреляций и соотношений и, можно сказать, буквально научился на этих языках говорить.

Теперь, когда вы предлагаете Google Translate перевести какой-то текст, он не просто вываливает на вас набор слов из бесчувственного словаря, а строит фразу, причём не только используя известные ему алгоритмы, но и параллельно совершенствуя их, самообучаясь.

«Получив предложение на английском языке и запрос перевода на немецкий, сервис сканирует все известные ему документы на английском и немецком, ищет точное совпадение и затем возвращает соответствующий текст на немецком языке. [...] Компьютеры оценивают статистические закономерности в больших массивах ранее накопленного цифрового контента, создание которого потребовало больших затрат, но воспроизведение их не стоит почти ничего».

ЭРИК БРИНЬОЯФСОН, директор Центра цифрового бизнеса MIT,
ЭНДРЮ МАКАФИ, социолог и экономист из MIT

Является ли такого рода «интеллектуальная» работа собственно интеллектом? И да, и нет. Если достаточно грубо обобщить то, что делает какой-нибудь Google Translate, то ситуация выглядит следующим образом:

- на его условный сенсор приходит некая информация, которую он должен сначала как-то распознать, идентифицировать, а потом отреагировать на неё;
- чтобы придумать адекватную реакцию на этот идентифицированный им как-то раздражитель, он должен сопоставить её с каким-то своим опытом – с тем, что он уже знает;
- он обращается к своей базе данных (своему «опыту») и находит некое соответствие, и производит реакцию – некое поведение, а проще говоря, предлагает перевод слова или фразы.

Если вы имеете хоть какое-то представление о том, как развивалась психотерапевтическая наука, а слова «бихевиоризм» и «когнитивно-поведенческая психотерапия» для вас не пустой звук, то вы уже понимаете, что мы имеем дело с классической моделью психики.

ПСИХОТЕРАПЕВТИЧЕСКИЕ РАЗВИЛКИ

Когда говорят о психотерапии, то первым делом вспоминают, конечно, классический психоанализ. Фигура Зигмунда Фрейда и правда была грандиозной, а думать о сексе, согласитесь, куда приятнее и привычнее, чем о чём-либо ещё. Но психоанализ ни в практическом, ни в концептуальном плане не выдержал испытания временем. Это красивая теория, но она ненаучна.

Наука начинается с нейрофизиологии – с наших соотечественников И. М. Сеченова, И. П. Павлова, А. А. Ухтомского, Л. С. Выготского, А. Р. Лурии, П. К. Анохина... Об открытиях этих великих учёных я уже много раз рассказывал в других своих книгах, а здесь нас интересуют лишь общие схемы – на них и остановимся.

Конец ознакомительного фрагмента.

Текст предоставлен ООО «ЛитРес».

Прочитайте эту книгу целиком, [купив полную легальную версию](#) на ЛитРес.

Безопасно оплатить книгу можно банковской картой Visa, MasterCard, Maestro, со счета мобильного телефона, с платежного терминала, в салоне МТС или Связной, через PayPal, WebMoney, Яндекс.Деньги, QIWI Кошелек, бонусными картами или другим удобным Вам способом.