

«В 2050 году людям живется безумно тяжело... Роботы используют их как подсобную рабочую силу, которую содержат в специальных лагерях, разбросанных подобно ГУЛАГу по всей территории Земли. Там тусклый искусственный свет, там едва топят -- лишь для того, чтобы люди не погибли от холода, и, конечно, никаких признаков комфорта в нашем прежнем понимании». (Кевин Уорвик «Наступление машин»)

НЕСКОЛЬКО ПРАВДОПОДОБНЫХ СЦЕНАРИЕВ



Однажды маститого бельгийского профессора Юго де Гари пригласили в Японию для разработки Robokoneko -- искусственного робота-котенка. Работая над проектом, де Гари преисполнился пессимизма: он испугался того, что роботы, обладающие интеллектом, станут не помощниками, а могущественными врагами человечества. При этом ученый нисколько не сомневался, что люди в результате утратят свое доминирующее положение на Земле. «Надвигается царство искусственного интеллекта, или арtilекта, который в один прекрасный день может решить, что мы нечто вроде вредных насекомых и нас нужно затравить», -- написал де Гари в одной из своих статей.

Примерно в то же время изобретатель и мыслитель Рэй Курцвайль в книге «Эра одушевленных машин» высказал мысль, что для людей скоро просто не останется места на Земле. Машины уничтожат нас как абсолютно бесполезные и потенциально опасные существа.

Так, впервые о вреде интенсивного развития новых технологий заявили не «зеленые», не правозащитники и последователи далай-ламы, а люди, находящиеся на переднем плане науки. Дальше было еще страшнее.

Знаменитый Унабомбер, профессор математики Теодор Качински, предсказал, что человечество в скором будущем будет «выключено» из развития цивилизации. В лучшем случае люди превратятся в нечто вроде домашних животных при заполняющих мир роботах. В худшем -- Земля переживет стадию разобщенности (dystopia), когда нескончаемые войны между разумными машинами и последними представителями человеческой расы приведут к полному исчезновению homo sapience как биологического вида.

«Нет, -- возразил ему главный научный сотрудник Sun Microsystems Билл Джой в недавней истеричной статье, -- нас обязательно уничтожит какой-нибудь злой гений, который создаст целую армию микроскопических нанороботов. Мы уже проскочили точку технологического невозврата: бесконтрольное самовоспроизведение робо-, гено- или наномеханизмов вот-вот должно начаться».

Наконец, самый свежий пример: английский изобретатель Кевин Уорвик в своей не слишком заметной книге «Наступление машин» напроорочествовал, будто в недалеком будущем миром будут править роботы, чей интеллект намного превзойдет человеческий. Люди станут для них рабами.

«Человеческая раса, похоже, играет свою последнюю партию, -- пишет Уорвик. -- Период нашего господства на Земле подходит к концу. Мы можем надеяться лишь на то, что машины будут обращаться с нами так, как мы обращаемся с другими животными, -- они сделают нас рабами или поместят в зоопарки. Разве мы этого хотим?»

ОБ ОДНОМ НЕУДАВШЕМСЯ ЗАГОВОРЕ

Прежде всего привлекают внимание имена этих новоявленных нострадамусов. Перед нами крупнейшие ученые современности, авторитетнейшие специалисты в области высоких технологий и не просто кабинетные теоретики, но люди, непосредственно занимающиеся практическими исследованиями. То, где они работают, говорит само за себя: Брюссельский исследовательский институт Starlab, электронный гигант Sun Microsystems, Университет Рединга в Великобритании и т.д. Сегодня эти светила современной науки все как один похожи на испуганного школьника-отличника, смастерившего шутовину, которая вот-вот разнесет папин гараж. Они боятся. Боятся войны людей с разумными машинами. Они говорят об этом, как о свершившемся факте,

как о том, что обязательно будет. Обя-за-тель-но! Вопрос в том, кто победит, -- некоторые мечтатели все еще верят, будто победят люди, однако таких с каждым днем становится все меньше.

Дискуссии о том, что нам делать с нашим будущим, идут второе десятилетие и неизменно заходят в тупик полной роботизации планеты. Тот же Билл Джой упоминал о множестве ученых конференций, где постоянно обсуждаются данные вопросы. А в США существует президентская комиссия по будущему информационных технологий, периодически публикующая всевозможные неутешительные прогнозы.

Долгое время международное научное сообщество считало, что опасные разработки необходимо ограничить, то есть установить тотальный контроль за технологиями и создать этический кодекс для ученых, подобный, к примеру, клятве Гиппократы для врачей. Еще американский физик Лео Силард, один из создателей атомной бомбы, пытался убедить американских конгрессменов приостановить опасные разработки. Он был первым, кто предложил ученым всего мира компромисс между общечеловеческими ценностями и жадной прогресса: самостоятельно ограничивать некоторые области исследований.

Этот «тихий» заговор просуществовал 60 лет. Еще совсем недавно многие компании, вроде Sun Microsystems или Thinking Machines Corp., очень жестко контролировали собственные разработки в потенциально опасных областях. Чем еще можно объяснить то, что про чудеса электронных имплантантов мы узнаем только теперь, спустя 20 лет после начала их активного использования (в 80-х они применялись как эффективные сердечные стимуляторы). Не стал заканчивать свои опыты над человеческим мозгом знаменитый американский психолог Генри Меррей, проводивший исследования по заказу ЦРУ. Норберт Винер -- папа кибернетики -- отказался вплотную заниматься проблемой искусственного интеллекта, чего от него так хотели военные.

Из последних примеров: де Гари нарочно затягивал разработку мыслящего робота-котенка, вызвав недовольство компании-заказчика. Унабомбер вообще начал планомерно взрывать специалистов по новым технологиям, пока его не поймали и не бросили в одиночную камеру тюрьмы Сакраменто. Наконец, во множестве западных университетов стали готовить специалистов по био- и социоэтике, главная задача которых состояла как раз в том, чтобы определять потенциальную опасность для общества и человека тех или иных исследований.

Правда, все это время ученые никак не могли решиться на официальное введение повсеместного контроля за новейшими технологиями. И вдруг оказалось, что они опоздали. На научном небосклоне появилось маленькое облачко -- некий Кевин Уорвик, о котором мало кто знал.

ИСТОРИЯ ПАРАНОЙИ

Уорвика нельзя назвать крупным ученым. С начала семидесятых он занимался в British Telecom научной работой. Потом Кевин преподавал кибернетику в Имперском колледже науки, техники и медицины в Лондоне, Оксфорде, Ньюкасле и Университете Рединга. В последнем он, собственно говоря, и занялся детской забавой -- созданием всевозможных маленьких роботов и экспериментами с их электронными мозгами. В один солнечный день Уорвик чуть раньше обычного пришел в лабораторию, наводненную бегающими и ползающими механизмами, и застучал одного из своих любимцев за тем, что его просто ужаснуло: робот, не получив разрешения или стартовой программы, нашел через интернет другого робота, из Нью-Йорка, и самостоятельно передал ему информацию о том, как можно двигаться, не натываясь на вещи в комнате. Уорвик отключил шустрый механизм и стал думать. Мысль его растекалась примерно так: сегодня в мире существуют роботы с интеллектом насекомых, через пять лет их мозг будет равняться мозгу кошки, а через десять -- человека. Тогда машины начнут самостоятельно мыслить, потребуют право голоса и... война с ними станет неизбежна. Но главное -- в этой войне они обязательно победят.

Что же надо сделать для победы над разумными роботами? Все просто -- людям надо самим стать машинами. Или, по крайней мере, научиться управлять ими с помощью собственного разума. Сказано -- сделано. В августе 1998 года профессор нарушил молчание ученой общественности с помощью небольшой операции по имплантированию самому себе микрочипа. Этот чип преобразовывал биоэлектрические импульсы в радиоволны и транслировал их в компьютер. Но самое главное -- микросхема также была готова улавливать сигналы, посланные компьютером, превращая их в нервные импульсы! «Мы хотим перехватывать команды, отданные мозгом руке, и записывать их на жесткий диск, -- пояснял Уорвик журналистам. -- Тогда любые чувства -- радость, уныние или боль -- можно записать на жесткий диск компьютера и при надобности транслировать прямо в мозг».

Но как обыкновенный компьютер может читать человеческие мысли? Ведь мысль есть нечто мимолетное, неосоздаемое. Поэтому, чтобы уловить мысль, надо устремиться за

ней туда, где она возникает...

ИЗНАНКА МЕЧТЫ

Сейчас для измерения активности головного мозга используют обычную электроэнцефалограмму (ЭЭГ) -- биотоки регистрируют с помощью датчиков,



закрепленных на голове человека. Подобный метод безопасен, но весьма неточен, да и скорость его невысока. Поэтому лучше измерять электрические потенциалы непосредственно в головном мозге. Для этого имплантируют электроды внутрь черепа. Сделать это непросто, да и реакция иммунной системы будет бурной -- она попытается отторгнуть чужеродный элемент, и это может грозить человеку гибелью. Сегодня, помимо чипов, Уорвиком и его коллегами используются так называемые электроэнцефалограммные шлемы, чьи электроды плотно прилегают к голове человека. Они фиксируют характерные биотоки, возникающие при малейшем усилии мысли. Остается лишь ввести электромагнитные сигналы в компьютер. Там и будут храниться характеристики различных биотоков: их сила, форма, частота и другие параметры. По мере накопления данных исследователи построят компьютер, читающий мысли, он будет работать примерно так же, как машина, распознающая голоса: та анализирует звуковые волны и соотносит с ними слова. Точно так же отдельные биотоки фиксируют мысли или хотя бы направление, в котором думает человек. Казалось бы, просто. Увы, мыслить мы не так, как говорим, мыслить мы довольно хаотично.

Поэтому сразу после фиксации биотоков начинается главный этап работы: наведение порядка в хаосе роящихся мыслей. «Анализировать электроэнцефалограмму все равно, что сидеть во время шумной вечеринки за плотно закрытой дверью и пытаться понять, кто и что говорит, -- признается один из помощников Уорвика. -- Кое-что понимаешь: взрыв хохота, топот танцующих, но какие разговоры ведутся за всеми столами подряд -- не поймешь».

Из этого хаоса, как из знаков незнакомого языка, ученые вылавливают отдельные образчики биотоков. Если человек напряженно думал о чем-то, можно заставить компьютер выполнить это. В соответствии данному узору мысли ставится команда. Для

эффективной работы ученые должны ограничиваться двумя или тремя различными образчиками биотоков, выделяя их из какофонии мозга. Этого достаточно для простейших компьютерных операций. А для непростейших?.. Как печатать на компьютере с помощью мысли?

Ученые решили для этого воспользоваться «виртуальной клавиатурой», на которой можно было выбрать любые буквы, любые клавиши, а для этого -- проделать несколько нехитрых операций. Стоило, например, представить себе вращающийся куб, как мозг порождал определенный биоток, включавший клавишу «влево». Если человек вспоминал какую-то мелодию, которая активизировала другую часть мозга, то возникали совсем иные биотоки, управлявшие клавишей «вправо». И так далее. В конце концов, следуя подобной методе, человек мог выбрать любую букву.

НЕДАВНО УОРВИК ПОВТОРИЛ СВОЙ ЭКСПЕРИМЕНТ

Кевин, панически боящийся высоты, поднялся на крышу небоскреба. В плечо профессора был имплантирован чип, подключенный к нервным окончаниям. Все импульсы страха, пробежавшие по электронным нервам, тотчас преобразовывались в цифровую форму и транслировались в интернет. А в это время на другом берегу Атлантики, в Англии, эти биты и байты оцифрованных чувств поступали в чип, внедренный в тело жены Уорвика. И когда она вдруг испытала сильнейший страх именно высоты, эксперимент посчитали удавшимся: супруги Уорвик стали первой семейной парой, связавшей свою судьбу электронными узами.

Кевину тут же поступило множество предложений о сотрудничестве от крупнейших корпораций и правительств некоторых стран. Во всем мире активизировались разработки в области кибернетики и робототехники. На выставках, посвященных новым технологиям, объявилось множество «домашних» роботов, о создании которых раньше никто не сообщал. То есть Уорвик сделал своеобразную отмашку: мы не должны бояться делать машины, но мы должны знать, как их победить или как управлять ими. Система ограничений, которую филигранно выстраивали ученые разных стран, пошла прахом. «Вводить запрет на подобные работы бесполезно, -- заявил профессор. -- Всегда найдется какой-нибудь «гений», который втихомолку создаст армию всемогущих киберов, стремясь подчинить мир себе». Но что же предложил нам английский ученый взамен?

БУДУЩЕЕ ПРОФЕССОРА УОРВИКА

«Давайте все станем киборгами», -- предложил он. Микрочип размером с половинку спички, который Уорвик вживил себе в левую руку, -- всего лишь первый шаг к этому. Дальше ученый собирается модернизировать человека электронными модулями памяти, устройствами инфракрасного зрения, онлайнowymi системами управления удаленными объектами и другими высокотехнологичными прибамбасами. Однако Кевин явно не рассматривает все возможные сценарии, которые повлечет за собой его инициатива. Мы же с вами вполне можем предаться кибернетическим мечтам.

Итак, представьте себе мир, наводненный киборгами с кремниевыми телами. Представьте себе бесконечное электронное бессмертие в масштабах человечества. Вообразите, как детей сразу после рождения будут снабжать механическими руками-ногами и электронными чипами. Или как их будут делать такими на фабриках. Хлоп -- ребенок, хлоп -- другой. «Запускайте второй конвейер!»

Только подумайте, что исчезнет преступность, исчезнет вместе с гражданской свободой, ведь киборгов с вживленными чипами можно будет легко контролировать.

Наконец, взгляните на такую картину: киборги организуют митинги и демонстрации по всему миру, требуя для себя избирательных прав. Их пытаются разогнать люди-полицейские, но на полуроботов не действуют дубинки и слезоточивый газ. В результате под железным (в прямом смысле) давлением инициируется общемировой судебный процесс, вроде того, на котором решали, есть ли душа у индейцев. Люди после долгих препирательств отказывают киборгам в их требованиях. Начинается война. Полуроботы побеждают, и уже новые проблемы встают перед ними.

Где именно проходит грань между киборгами и людьми? Если у меня железный палец, я что, киборг? А если человеческой у меня осталась только голова?

Потом, что делать со всевозможными компьютерными вирусами? Уже сегодня производители «домашних» роботов бьют тревогу -- их игрушки можно заставить делать все что угодно, просто взломав программу. А представьте, как у вас вдруг без команды оживает механическая рука. Возьмет да и задушит своего хозяина...

И справится ли человеческая психика с особенностями своего кибернетического тела? Не приведет ли несовместимость органики с электроникой к самоуничтожению киборгов, которые просто-напросто сойдут с ума?

Молчит Кибальчиш технического прогресса. Не дает ответа. По его понятиям, нас ждет незавидное будущее кремниевых полукровок, неведомых зверюшек, утративших человеческие черты и взамен приобретших... а что, собственно говоря, мы получим взамен? Механические пальцы? Чипы под кожей, чтобы открывать двери на расстоянии? Может, остальные ученые с их «тихим» заговором были столь же не правы, но их действия по крайней мере не вели к таким, мягко говоря, неоднозначным последствиям.

Потому есть предложение. Пока еще не поздно, пока английский доктор Но в результате своих экспериментов не успел создать того, против чего он сам собирается бороться, а новый завод, построенный с вами по соседству, еще не выпускает киборгов пачками, Кевина Уорвика надо пристрелить. Ради будущего человечества.

Кирилл ЖУРЕНКОВ, Александр ВОЛКОВ

В материале использованы фотографии: Camera PRESS/FOTOBANK

<http://www.ogoniok.com/archive/2002/4747-4748/19-36-37/>