

ЭЛЕКТРОНИКА

НАУЧНО-ПРАКТИЧЕСКОЕ ИЗДАНИЕ

№ 2 | март-май | 2022

ПЛЮС

ТЕМА НОМЕРА:

ИСКУССТВЕННЫЙ ИНТЕЛЛЕКТ



E-mail: sym@riftek.com
Тел.: +375 17 281 36 57



**РИФТЭК
SMT**
АВТОМАТИЧЕСКИЙ МОНТАЖ ПЕЧАТНЫХ ПЛАТ

ЧУП «РИФТЭК-СМТ»
Республика Беларусь,
220090, г. Минск,
Логойский тракт, 22
УНП 192241841

БелСканту

TEKO INSTRUMENTS	COMET	ZEHO	SAURIS
SOCIOPEX	COMET	Penata	SAURIS
HEALTER	COMET	PROCELL	SECO
KDS	COMET	Neoway	INTEP
MC	COMET	AO	SECO
4 SINPRO	COMET	AO	SECO

ООО «БелСКАНТИ»
+375 (17) 256-08-67, 398-21-62
nab@scanti.ru
www.scanti.com
Стр. 64
УНП 190813939



intel
XEON
inside™

ANALOG DEVICES
Honeywell
Hittite
SICK

ТУП «Альфачип Лимитед»

Поставка электронных компонентов,
средств автоматизации, компонентов
для светодиодного освещения

220012, г. Минск, ул. Сурганова, 5а, 1-й этаж
Тел./факс: +375 17 366 76 01, +375 17 366 76 16
факс: +375 17 366 78 15
www.alfa-chip.com
www.alfacomponent.com
УНП 192525135



TOSHIBA

toshiba.com

**Относитесь
к будущему
с пониманием!**

ИЗДАЕТСЯ ПРИ ИНФОРМАЦИОННОЙ ПОДДЕРЖКЕ ФАКУЛЬТЕТА РАДИОФИЗИКИ
И КОМПЬЮТЕРНЫХ ТЕХНОЛОГИЙ БЕЛОРУССКОГО ГОСУДАРСТВЕННОГО УНИВЕРСИТЕТА

НОВОСТИ

ХОЛОДНАЯ ВЕСНА И ГОРЯЧИЕ НОВОСТИ... 2

ИСТОРИЯ

ПРИКЛЮЧЕНИЯ МИКРОПРОЦЕССОРА В СССР 6

ТЕМА НОМЕРА

ПЯТЬ ПРЕДСКАЗАНИЙ IBM О ТОМ, КАКОЕ БУДУЩЕЕ НАС ЖДЕТ В 2022 ГОДУ 14

ИСКУССТВЕННЫЙ ИНТЕЛЛЕКТ. ВСЕ, ЧТО ВЫ ХОТЕЛИ ЗНАТЬ 16

НАУКА

РАСПОЗНАВАНИЕ ПОЛНОЦВЕТНЫХ РАСТРОВЫХ ИЗОБРАЖЕНИЙ
ПОЛЯРНЫХ СИЯНИЙ С ПОМОЩЬЮ НЕЙРОННЫХ СЕТЕЙ ГЛУБОКОГО ОБУЧЕНИЯ
Коберник-Березовский Я.А., Чекановский П.А., Кочетова Д.А.,
Светашев А.Г., Турышев Л.Н. 39

ИССЛЕДОВАНИЕ ВОЗМОЖНОСТИ ОЦЕНКИ КАЧЕСТВА МОТОРНОГО МАСЛА
ПО ЭЛЕКТРОФИЗИЧЕСКИМ ПАРАМЕТРАМ ПРИ ЭКСПЛУАТАЦИИ АВТОМОБИЛЕЙ
С.С. Беленькая, В.С. Курило, Е.С. Максимович 46

ЛАЗЕРНОЕ ВОЗДЕЙСТВИЕ НА ТКАНЬ 07C11-KB
С КОМБИНИРОВАННЫМ ПОКРЫТИЕМ ЦИРКОНИЯ И УГЛЕРОДА
М.И. Маркевич, В.И. Журавлева, Т.Ж. Кодиров, И.П.Акула,
Н.М. Чекан, Е.Н. Щербакова 50

МНОГОПОЛЬЗОВАТЕЛЬСКАЯ СИСТЕМА УДАЛЕННОЙ 3D FDM ПЕЧАТИ
Б.С. Петкевич, И.А. Романов 53

ПРАЙС-ЛИСТ 59

ЭЛЕКТРОНИКА ПЛЮС ЦИФУС

**№2
март-май 2022**

Издание для специалистов, занимающихся разработкой и поставкой электроники, компонентов и другой продукции в различных отраслях промышленности. Издание знакомит специалистов с новыми достижениями и разработками в области электроники, микроэлектроники, электротехники, оптоэлектроники, энергетики, средств связи. Публикует научные статьи ученых. Размещает рекламу по теме номера.



Учредитель:

ООО «ВитПостер»

Главный редактор

Бокач Павел Викторович
m6@tut.by
+375 (29) 338-60-31

Секретарь:

Садов Василий Сергеевич, к.т.н.
sadov@bsu.by

Подписано в печать 20.06.2022.

Отпечатано в типографии
ООО "ЮСТМАЖ",
ул. Калиновского, 6 Г 4/К,
220103, г. Минск
ЛП №02330/250

Бумага офсетная.
Тираж 299 экз. Заказ 198.

Издатель ООО «ВитПостер».
Свидетельство о государственной регистрации
издателя, изготовителя, распространителя
печатных изданий № 1/99 от 02.12.2013.
E-mail: artmanager3@mail.ru

© ООО «ВитПостер», 2022

ХОЛОДНАЯ ВЕСНА И ГОРЯЧИЕ НОВОСТИ...

INTEL ПОДТВЕРДИЛА ЗАДЕРЖКИ В ВЫПУСКЕ ВИДЕОКАРТ ARC

Вице-президент Intel Лиза Пирс сообщила, что мобильные видеокарты Intel Arc официально выпущены в продажу, однако найти их очень сложно, и виной тому срыв цепочек поставки комплектующих. Компания обещает, что серия мобильных карт Arc 3 появится как можно скорее, а

серии Arc 5 и Arc 7 «в начале лета». Что касается настольной версии, то вначале эти ускорители будут доступны для OEM. Настольные карты начального уровня серии A сначала появятся в Китае, причём по-прежнему планируется второй квартал. Для китайского рынка выпуск

будет осуществлён как для розничной продажи, так и для OEM. Глобальная доступность этих карт начального уровня ожидается немного позднее. Что касается остальной линейки, Arc 5 и Arc 7, то она появится в продаже «позднее этим летом».

intel.com

INTEL: ДЕФИЦИТ МИКРОСХЕМ ПРОДЛИТСЯ ДО 2024 ГОДА

Многие аналитики выражали надежду, что ситуация с дефицитом микросхем улучшится в 2022 году. Однако до сих пор полупроводниковые компании не имеют возможности заполнить склады. И вот теперь исполнительный директор Intel Пэт Гелсингер сообщил, что ожидает продолжения дефицита до 2024 года.

Пэт Гелсингер уверен, что дефицит продлится намного дольше, чем ожидалось, из-за нехватки специальных производственных инструментов, что не позволяет нарастить объёмы производства.

Цены на видеокарты продолжают снижаться, однако видеокарты – это

далеко не весь рынок микросхем, а лишь небольшая часть этой индустрии. Другие компании-производители бытовой электроники, автомобильные компании по-прежнему страдают от дефицита чипов.

В сложившейся ситуации Intel находится в более выгодном положении благодаря фирменной стратегии IDM 2.0. В то время, как большинство производителей связаны с поставщиком TSMC, у Intel есть собственные производственные мощности. Кроме того, при необходимости, компания может заказывать микросхемы и сторонних компаний, что обеспечивает лучшую доступность продукции Intel на рынке.

intel.com



AMD ГОТОВИТ ПРОЦЕССОР ДЛЯ ШЛЕМА ВИРТУАЛЬНОЙ РЕАЛЬНОСТИ



Компания Magic Leap готовит шлем дополненной реальности, который может быть основан на специализированном процессоре от AMD. В тесте Basemark обнаружена SoC от AMD, которая может быть использована в шлеме с кодовым именем Mero. Эта SoC объединяет процессор на базе архитектуры Zen 2 и графику на основе RDNA2. Процессор имеет 8 ядер, однако может быть, что речь идёт о 4 ядрах и 8 потоках. В настоящее время количество вычислительных блоков RDNA2 в iGPU неизвестно. В качестве операционной системы Magic Leap использует Android 10 для архитектуры x86-64. Шлем оснащён дисплеем 720×920 пикс., а объём памяти, доступной для ОС, составляет 1 Гб (не считая памяти, выделенной для GPU).

nvworld.ru

AMD ПРЕДСТАВИЛА RYZEN 7000

В ходе Computex 2022 компания AMD официально представила серию процессоров Ryzen 7000. Новый процессор основан на архитектуре Zen 4, которая изготавливается по 5 нм нормам. Чип имеет вдвое больший объём кэша L2 и более высокую тактовую частоту, до 5,5 ГГц в реальной работе с играми. Процессор обеспечивает 15% прирост производительности в одном потоке. Кроме того, AMD продемонстрировала работу CPU в многопоточной задаче Blender, где рендер на топовой модели AMD оказался на 30% быстрее, чем на Intel Core i9 12900K. В дополнение, серия Ryzen 7000 содержит совершенно новое ядро ввода-вывода, изготовленное по 6 нм процессу, графику AMD RDNA 2, новую архитектуру со сниженным энергопотреблением, заимствованную из мобильных процессоров Ryzen, а также ряд новых технологий связи (DDR5, PCI Express 5.0 и т. д.) и поддержку до 4 дисплеев. AMD вступает в более тесную борьбу с Intel в высшем ценовом сегменте и надеется ещё сильнее закрепиться на этом рынке.



nvworld.ru

SEAGATE И PHISON БУДУТ ВМЕСТЕ РАЗРАБАТЫВАТЬ ТВЕРДОТЕЛЬНЫЕ НАКОПИТЕЛИ КОРПОРАТИВНОГО КЛАССА

Компания Seagate Technology, специализирующаяся на решениях для инфраструктуры хранения больших объемов данных, и компания Phison Electronics, известный разработчик контроллеров для твердотельных накопителей, объявили, что планируют расширить свои портфели твердотельных накопителей за счет нового поколения высокопроизводительных корпоративных твердотельных накопителей высокой плотности с поддержкой NVMe. Новые накопители помогут снизить совокупную стоимость владения (TCO) за счет увеличения плотности хранения, снижения энергопотребления и повышения производительности. Стороны объявили о заключении долгосрочного партнерства, охватывающего циклы разработки и распространения твердотельных накопителей корпоративного класса.



Seagate и Phison сотрудничают с 2017 года, контроллеры Phison были выбраны для SSD Seagate массового сегмента с интерфейсом SATA. Сотрудничество распространилось на линейку высокопроизводительных потребительских твердотельных накопителей NVMe PCIe Gen4x4 FireCuda и первых в

мире специально разработанных твердотельных накопителей NAS NVMe. Сотрудничество удовлетворит растущий глобальный корпоративного спроса на более высокую плотность, более быструю и интеллектуальную инфраструктуру хранения.

ixbt.com

ВНЕШНИЙ SSD KINGSTON С АППАРАТНЫМ ШИФРОВАНИЕМ И СЕНСОРНЫМ ЭКРАНОМ

Компания Kingston представила линейку SSD-накопителей IronKey Vault Privacy 80. Это первый для компании ОС-независимый внешний SSD с сенсорным дисплеем и аппаратным шифрованием данных. Kingston IronKey Vault Privacy 80 совместим с Microsoft Windows, macOS, Linux, Chrome OS и любой другой операционной системой, которая поддерживает работу с USB-накопителями. В линейку входят модели емкостью 480 ГБ, 960 ГБ и 1920 ГБ, производитель предлагает три года гарантии и бесплатного технического обслуживания.

IronKey Vault Privacy 80 сертифицирован по американскому стандарту FIPS 197 с 256-битным ключом XTS-

AES и способен защитить данные от брутфорс-атак и взлома BadUSB. Можно задать несколько паролей (администратор/пользователь), выбор между PIN и парольной фразой, настройка правил ввода пароля (длина от 6 до 64 символов, буквенно-цифровые сочетания). Устройство автоматически заблокируется, если не будет использоваться в течение заданного периода времени, а если пароль будет неправильно введен 15 раз подряд, то произойдет криптостирание данных (при этом удаляются все ключи безопасности). Kingston IronKey Vault Privacy 80 может работать в режиме защиты от записи, чтобы исключить попадание вредоносного ПО. При всем этом, накопи-



тель вполне компактен: его размеры составляют 122,5x84,2x18,4 мм.

kingston.com

КАКИЕ СМАРТФОНЫ БОЛЬШЕ ВСЕХ ИЗЛУЧАЮТ

Данные BanklessTimes позволяют узнать, какие из современных смартфонов имеют наибольший уровень электромагнитного излучения. Опубликован список моделей с максимальным уровнем электромагнитного излучения. Самым худшим оказался Motorola Edge с показате-

лем SAR 1,79 Вт/кг. Проблема в том, что это почти на 0,2 Вт/кг, нежели максимальное значение, рекомендуемое Федеральной комиссией по связи США (1,6 Вт/кг). И это единственный смартфон, который вышел за пределы данного лимита. Ещё парочка аппаратов – ZTE Axon 11 5G и

OnePlus 6T – максимально приблизились к лимиту. Остальные лидеры антирейтинга имеют ощутимо меньший показатель SAR. При этом стоит отметить, что в десятке лидеров три смартфона Google, парочка OnePlus и две модели Sony.

BanklessTimes

УСТРОЙСТВО MAGEWELL USB FUSION ДЛЯ ВИДЕОЗАХВАТА

Компания Magewell представила устройство видеозахвата USB Fusion с несколькими входами, функциями коммутации и объединения видеопотоков. Оно предназначено для проведения презентаций, онлайн-лекций, вебинаров, прямых трансляций, видеоконференций, новостных репортажей и других подобных применений.

Имея два входа HDMI 1.4 и один вход USB 3.0 для веб-камеры, USB Fusion позволяет пользователям переключаться между источниками HD 1080p60 или объединять два входа в один выход (PiP или PnP) HDMI 1.4 и передавать результат в популярное программное обеспечение через порт USB 3.0 Type-C. Кроме того, есть два

разъема диаметром 3,5 мм – линейный вход и выход на наушники.

Кнопки на устройстве позволяют переключаться между полноэкранными источниками или выбирать комбинированный макет сцены, а веб-интерфейс на основе браузера предлагает переключение сцен, регулировку громкости и настройку устройства. Полностью возможности USB Fusion раскрываются с помощью бесплатного приложения USB Fusion для планшетов iOS и Android. Интуитивно понятное приложение предоставляет расширенные элементы управления макетом для объединения любых двух из трех источников ввода и позволяет пользователям добавлять в презентации изображения, видео-

клипы, фоновую музыку и рукописные заметки. Докладчики могут заранее создавать списки воспроизведения, включающие живое видео и другие медиафайлы, а также записывать или делать скриншоты комбинированного вывода во встроенную память USB Fusion. Аннотации также можно экспортировать в файлы для архивирования или последующего обмена.

USB Fusion поддерживает ОС Windows, Mac и Linux. Соответствие спецификациям UVC и UAC обеспечивает широкую совместимость с программным обеспечением и простую установку без драйверов. Размеры устройства равны 186 x 101,2 x 30 мм. Цену компания не называет.

magewell.com

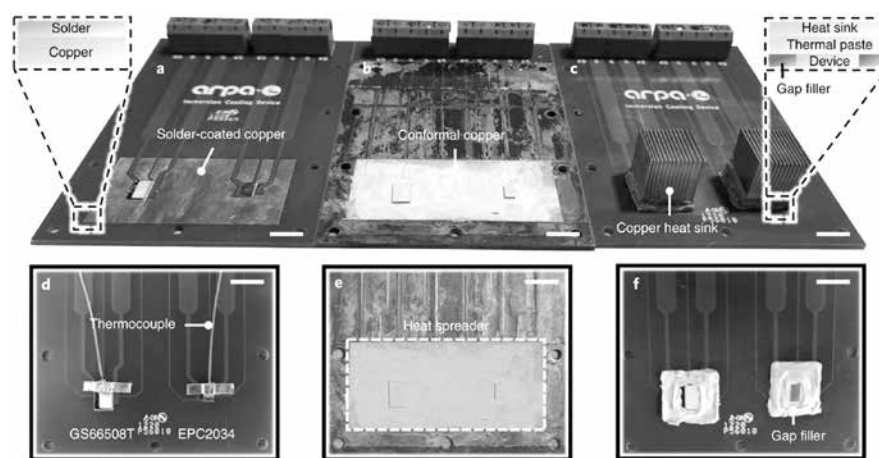


ИНЖЕНЕРЫ ПРИДУМАЛИ, КАК ПОВЫСИТЬ ЭФФЕКТИВНОСТЬ ОХЛАЖДЕНИЯ ПРОЦЕССОРОВ НА 700%

Инженеры из США придумали новый способ охлаждения полупроводниковых изделий, который не предполагает применение термоинтерфейса. Конформное покрытие в лабораторных условиях продемонстрировало свою высокую эффективность в сравнении с традиционными способами отведения тепла.

Электронные устройства нагреваются во время работы, и чем больше мощность, тем, при прочих равных, сильнее нагрев. Особенно остро эта проблема встает при миниатюризации устройств. Исследователи Иллинойский университет в Урбане-Шампейне придумали новый эффективный способ отведения тепла от полупроводниковых изделий.

Обычно к наиболее горячим местам микросхемы подсоединяют радиатор через термоинтерфейс. Например, термопасту. Вместо этого американские ученые предложили покрыть все устройство тонким слоем полимерного изолятора, а сверху «за-



ливать» его слоем меди, полностью повторяя форму устройства. Такое конформное покрытие плотно контактирует с устройством и эффективно отводит от него тепло.

Инженеры придумали способ повысить эффективность охлаждения процессоров на 700%

«Мы сравнили наше покрытие со стандартными методами теплоотвода.

Выяснилось, наш метод обеспечивает рассеивание на 740% больше мощности на единицу объема», – пояснил один из исследователей.

Ожидается, что в будущем это изобретение позволит сделать электронные устройства, в том числе телефоны и компьютеры, более мощными и компактными.

Газета.ru

В КИТАЕ ЭЛЕКТРОМОБИЛИ СРАВНЯЛИСЬ ПО ЦЕНЕ С ОБЫЧНЫМИ МАШИНАМИ

В Китае сосредоточено более половины мирового производства электромобилей. Ещё несколько лет назад о неизбежном удешевлении электромобилей говорили энтузиасты, но скептики упорно отказывались верить в это. Сегодня же электромобили почти сравнялись по цене с обычными машинами, правда, пока только в Китае.

По данным международного энергетического агентства (МЭА), средняя цена нового электромобиля в Китае примерно на 10% превышает среднюю цены похожего автомобиля на двигателе внутреннего сгорания. Чтобы ощутить важность этого достижения, нужно добавить, что в США и Европе разница в стоимости между

электромобилями и машинами с двигателями внутреннего сгорания достигает 50%, даже с учётом государственных скидок. При этом на многие премиальные модели никаких скидок не предусмотрено.

В Китае сосредоточено более половины мирового производства электромобилей, при этом по итогам 2021 году продажи электромобилей в стране выросли втрое. Также в Китае производятся батареи и основные комплектующие для таких машин. Ранее депутаты комитета по окружающей среде Европарламента поддержали план ЕС запретить продажи новых автомобилей с бензиновыми и дизельными двигателями с 2035 года. Новые правила должны исполнять 27 стран. Купленные ранее машины с двигателями внутреннего сгорания можно будет эксплуатировать как минимум до 2050 года.

китайские-автомобили.рф



РОССИЙСКИЙ ЭЛЕКТРОПОЕЗД «ИВОЛГА 3.0», СПОСОБНЫЙ РАЗГОНЯТЬСЯ ДО 160 КМ/Ч

«Трансмашхолдинг» заявил о том, что на Тверском вагоностроительном заводе завершилась приемка новейшего российского электропоезда «Иволга 3.0». Это позволяет начать производство партии поездов.

«На Тверском вагоностроительном заводе (ТВЗ, входит в состав АО «Трансмашхолдинг») успешно завершилась работа комиссии, которая оценивала результаты работы по созданию новейшего российского электропоезда постоянного тока ЭГЭ2Тв «Иволга 3.0», — отмечается в сообщении. По итогам работы приемочной комиссии был подписан акт, подтверждающий соответствие поезда техническому заданию. Он позволит осуществить производство поездов в 35 одиннадцативагонных составов.



Электропоезд «Иволга 3.0» создан российскими конструкторами на отечественной компонентной базе. Это первый электропоезд с конструкционной скоростью 160 км/ч, который полностью спроектирован и изготовлен в

России. Поезд имеет широкие (1400 мм) двери, что позволяет ускорить пассажирообмен на 15% в сравнении с представленными на российском рынке аналогами.

ТАСС

ЦЕНЫ НА ВИДЕОКАРТЫ РУХНУЛИ ВСЛЕД ЗА РЫНКОМ КРИПТОВАЛЮТ

Снижение стоимости биткоина и Ethereum на 30–40 % всего за одну неделю вызвало настоящий хаос на рынке видеокарт. Майнеры стали распродавать свои графические ускорители по заниженным ценам, поскольку доходность майнинга значительно снизилась. Это в свою очередь соз-

дало проблемы для производителей видеокарт, успевших накопить запасы моделей GeForce RTX 30-й серии и Radeon RX 6000, и теперь испытывающих трудности с их продажей.

Рынок подержанных ускорителей уже предлагает карты значительно дешевле рекомендованного ценника.

За минувшие две недели графические карты рухнули в цене на 10 %.

Некоторые энтузиасты продают свои запасы видеокарт сразу комплектами. Например, за 2500 долларов можно приобрести шесть подержанных GeForce RTX 3080.

VideoCardz

ПРИКЛЮЧЕНИЯ МИКРОПРОЦЕССОРА В СССР



Мы живем в удивительное время: компьютеры окружают нас со всех сторон. Любимый смартфон, ноутбук на работе, медицинские приборы, браслеты и часы. Умные рекламные табло, самокаты и автомобили. В основе каждого такого устройства лежит тот или иной микропроцессор. А простой микрокомпьютер размером со спичечный коробок (на базе Atmega или STM32) можно положить в карман или установить в качестве дверного звонка. Мы живем в будущем, не особенно-то его замечая. Но до начала 1980-х ни один советский радиолюбитель даже мечтать не мог о домашнем персональном компьютере.

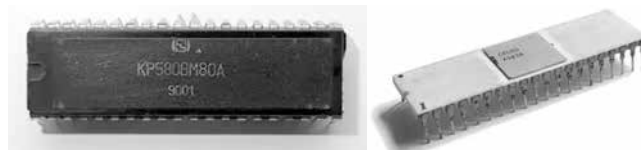
Сегодня мы попробуем взглянуть на первые микропроцессоры, применявшиеся на территории СССР. Статья целиком посвящена 8-битным CPU, которые так или иначе выпускались на территории Советского союза и СНГ.

Перед тем, как мы начнем разговор непосредственно о процессорах – не забывайте, что подавляющее большинство устройств, упомянутых в статье, закончили свой официальный жизненный цикл в лучшем случае в середине 1990-х годов. Из-за этого некоторые данные о них могут быть неточны и противоречивы. А о некоторых проектах, преимущественно военной направленности, практически нет никакой информации. Поэтому сосредоточимся мы преимущественно вокруг «гражданского» применения микропроцессоров на территории нашей страны. Первыми в очереди «на разбор» у нас 8-битные процессоры и устройства на их базе.

КР580ВМ80А

Пожалуй, не найдется ни одного радиолюбителя старше 30-35 лет, который бы не слышал об этом ле-

гендарном микропроцессоре. Фактически ВМ80А – это клон американского процессора Intel 8080А 1974 года. Выпускался с 1977 года (оригинал – с 1974). Первое название на отечественном рынке – КР580ИК80, но в ходе масштабного изменения системы обозначений микросхем в СССР получил свое «привычное» имя.



Отличия от своего американского собрата у процессора минимальны.

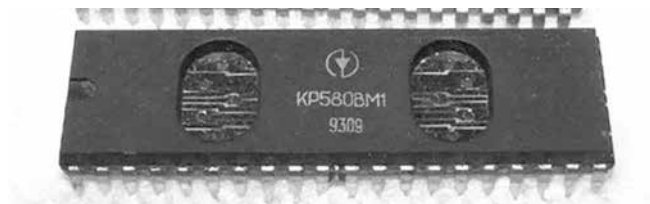
Кратко о технических характеристиках:

- тактовая частота – 2,5 МГц (максимальная гарантированная, но не предельная, если исходить из практики);
- самый распространенный формат – DIP40, но существовала и ранняя планарная версия с 48 контактами;
- шина адреса – 16 бит, до 64кБ оперативной памяти;
- шина данных – 8 бит;

- содержит 80 инструкций.

Разработан этот процессор был Киевским НИИ микроприборов под руководством А.В. Кобылинского. В поддержку процессора также был выпущен полноценный комплект дополнительных микросхем-аналогов серии Intel 82xx, среди которых КР580ВГ75 (контроллер дисплея), КР580ВН53 (таймер-счетчик, который также использовался в качестве музыкального «сопроцессора»).

В силу ряда исторических причин самостоятельного развития на территории нашей страны этот процессор не получил. Не имеющий аналогов в мире чип КР580ВМ1, выпущенный заводом «Квазар», содержал некоторые существенные улучшения, однако производился в ограниченном количестве.

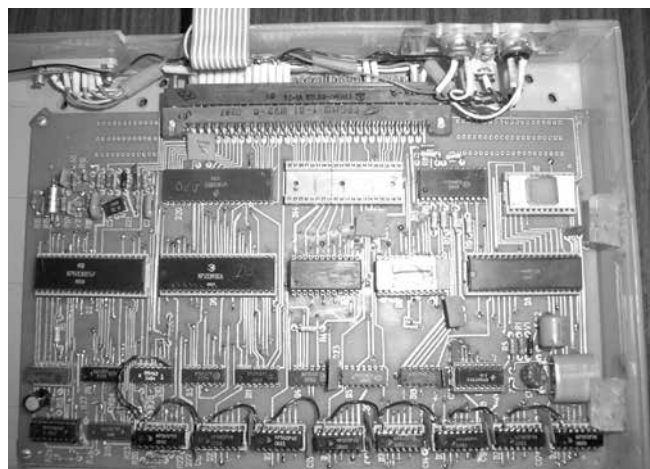


Сейчас этот процессор встречается лишь на сайтах для коллекционеров и стоит весьма ощутимых денег.

Теперь же обратимся к двум популярным компьютерам, построенным на базе процессора КР580ВМ80А.

«Радио-86РК»

Схема этого памятного для многих радиолюбителей компьютера впервые была опубликована в 1986 году, в журнале «Радио» №4-6. Авторами цикла статей числятся Д. Горшков, Г. Зеленко, Ю. Озеров, С. Попов. «86РК» (или «РК-86», как его иногда называют) позиционировался в качестве самодельного устройства, собрать которое в состоянии даже подросток, хотя бы раз державший в руках паяльник. Правда, добыть некоторые запчасти порой было весьма непросто.



Классическое исполнение самодельного «Радио-86РК»

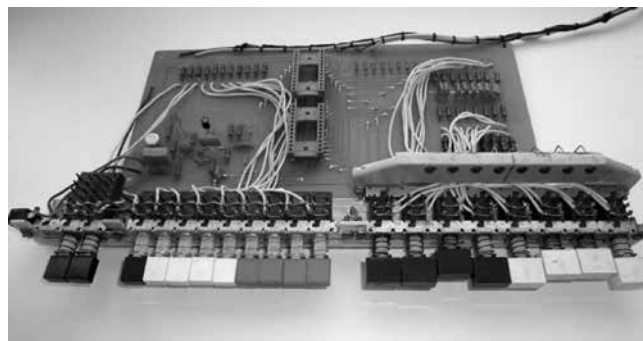
К слову, «Радио-86РК» является прямым потомком компьютера «Микро-80». Сам же «Микро-80» ввиду сложности сборки и большого (ок. 200 против 29 в «РК») количества микросхем популярности не снискал.



«Микро-80» и самодельная клавиатура

Характерная особенность компьютера – отсутствие графического режима и цветовой палитры. Однако читатели буквально заваливали редакцию «Радио» письмами с предложениями доработок и улучшений компьютера. Самые полезные и удачные идеи публиковались на страницах журнала в виде описаний и даже схем.

Классический «86РК» имел на борту 16-32кб оперативной памяти, одноканальный бипер-пищалку и два ПЗУ: первая – с программой «Монитор», вторая – с набором символов. Чтобы запрограммировать компьютер, сборщику требовался «ручной» программатор: микросхема ПЗУ вставлялась в панельку, а человек методично вводил данные с помощью кнопок, сверяясь с журналом. В котором, к слову, далеко не всегда печатались корректные прошивки. А если ошибка была допущена в процессе работы, микросхему приходилось стирать с помощью УФ-лампы.

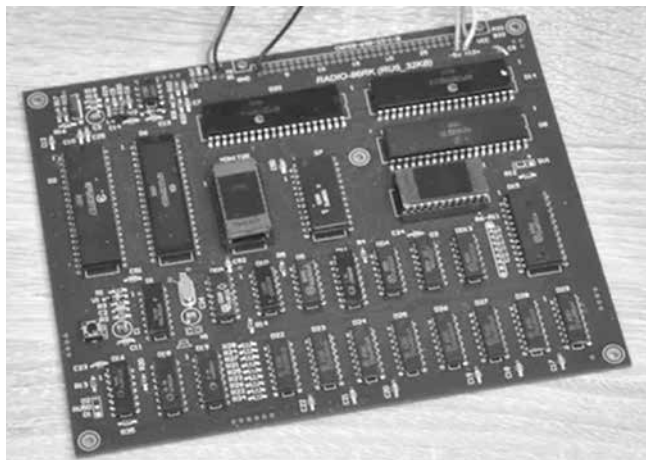


Один из вариантов программатора

Загрузка программ в компьютер осуществлялась через магнитофонный вход и (при внесении ряда доработок) через дискковод. Соответствующие модули расширения производились непосредственно авторами компьютера, и их можно было купить прямо в редакции «Радио» или у фирм-посредников.

Несмотря на массу ограничений и лишений, «РК» был полноценным компьютером: на нем можно было программировать, «прошивать» микросхемы с помощью самосборного программатора, играть в игры, слушать музыку – словом, делать всё то, к чему мы привыкли на

наших современных машинах. Конечно же, для полноценной «офисной» работы компьютер не годился. Но и не для предприятий он делался, так что здесь все в порядке.



Современная реализация компьютера, на 99% соответствующая оригинальной схеме

Существовала также коммерческая «доработка компьютера до цвета». Однако компания, производившая софт (преимущественно игровой) для этой версии компьютера, в середине-конце 1990-х годов разорилась, унеся с собой около 3-х десятков уникальных программ. Найти их сейчас не представляется возможным.

Современные клоны «Радио-86РК» «научились» работать в цвете, проигрывать музыку на популярном музыкальном чипе AY-3-8910 и «грузиться» с SD-карты или HDD.

К 1987 году началось промышленное производство «Радио-86РК» и его более продвинутых клонов: «Микро-ши», «Апогея», «Партнера», «Спектра» и прочих. Кроме того, в СССР были разработаны также 2 производных ПК: «Юниор ФВ-6506» и довольно-таки продвинутый «Электроника КР-04».

Прочие компьютеры на базе КР580ВМ80А, конструктивно превосходящие или принципиально отличающиеся от «Радио-86РК»:

- Текстовый компьютер «ЮТ88» (1989, схема опубликована в журнале «Левша»);
- «Орион» и его модификации;
- «Специалист».

А мы будем двигаться дальше. И на очереди у нас едва ли не самый интересный и самобытный компьютер советской эпохи.

«Вектор-06Ц»

История рождения этого интересного и по-настоящему самобытного компьютера весьма непростая. Он прошел путь от наброска на салфетке до промышленного производства и едва не был похоронен бюрократическими органами. Наиболее полную и интересную версию истории «Вектора» вы можете узнать на YouTube-канале главного «летописца» Вектора-06Ц Lafromm31. Неплохое текстовое изложение истории доступно здесь.



Одна из версий «Вектора» с характерными заводскими кабелями

Компьютер был разработан в Кишиневе в середине 1980-х годов. Авторство проекта принадлежит двум инженерам ПО «Счетмаш», Донату Темиразову и Александру Соколову.

КОМПЬЮТЕР ОБЛАДАЕТ СЛЕДУЮЩИМИ ХАРАКТЕРИСТИКАМИ:	
- СИСТЕМА КОМАНД МИКРОПРОЦЕССОРА	- КР580ВМ80А
- ЭФФЕКТИВНАЯ ТАКТОВАЯ ЧАСТОТА	- 2,4 МГц
- БЫСТРОДЕЙСТВИЕ	- 600 ТМС. ОПЕР./СЕК
- ОБЪЕМ ОПЕРАТИВНОЙ ПАМЯТИ	- 64 КБАЙТ
- ОБЪЕМ ДОПОЛНИТЕЛЬНОЙ ПАМЯТИ	- ДО 512 КБАЙТ
- "ЭЛЕКТРОННЫЙ КВАЗИ-ДИСК"	
- ГРАФИЧЕСКИЙ ДИСПЛЕЙ НА ВЫТВОРН ТЕЛЕВИЗИОННОМ ПРИЕМНИКЕ ЦВЕТНОГО ИЛИ ЧЕРНО-БЕЛОГО ИЗОБРАЖЕНИЯ	
- РЕЖИМ МНОХОХРОМНОГО (16 ВАРИАНТОВ ОКРАШИВАНИЯ - 1 ЦВЕТ+ЦВЕТ ФОНА)	
ОТображение ГРАФИЧЕСКОЙ ИНФОРМАЦИИ	- 512 X 256 ТОЧЕК (X, Y)
--/-- СИМВОЛЬНОЙ	- 42 X 25 (ЗНАКОМЕСТ, СТРОК)
--/--	- 64 X 32
--/--	- 85 X 25
- РЕЖИМ ЦВЕТНОЙ (15 ЦВЕТОВ И 1 ИЗ 16 ЦВЕТОВ ФОНА)	
ОТображение ГРАФИЧЕСКОЙ ИНФОРМАЦИИ	- 256 X 256 ТОЧЕК
--/-- СИМВОЛЬНОЙ	- 42 X 25 (ЗНАКОМЕСТ, СТРОК)
- ВНЕШНЯЯ ПАМЯТЬ НА ВЫТВОРН КАССЕТНОЙ МАГНИТОЛЕНТЫ	- 780 КБАЙТ/КАССЕТУ
- КЛАВИАТУРА РАБОТАЕТ ПО ПЕРЕКЛЮЧЕНИЮ В КОДАХ КОИ-7, КОИ-8	
- ЧАСЫ РЕАЛЬНОГО ВРЕМЕНИ С ДИСКРЕТНОСТЬЮ	- 0,02 СЕК
- ПРОГРАММНОЕ ОБЕСПЕЧЕНИЕ (ПРИ ОЗУ=60К):	
- МОНИТОР, АССЕМБЛЕР, РЕДАКТОР ТЕКСТОВ, МАКРОАССЕМБЛЕР, БЕЙСИК-ИНТЕРПРЕТАТОР	
- ПРОГРАММНОЕ ОБЕСПЕЧЕНИЕ (ПРИ ОЗУ=256К):	
- ОПЕРАЦИОННАЯ СИСТЕМА ОС/ВЕКТОР СОВМЕСТИМА С ОС СД/ИМ (DIGITAL RESEARCH, SDA)	
- ЯЗЫКИ ПРОГРАММИРОВАНИЯ: МАКРОАССЕМБЛЕР, БЕЙСИК, ПАСКАЛЬ, ФОРТРАН, ПЛ/М, ПЛ/1, С++	
СИСТЕМНАЯ ШИНА ОБЕСПЕЧИВАЕТ СОПРЯЖЕНИЕ С МОДУЛЯМИ РАСШИРЕНИЯ (МОНИТОР, ПРINTER И ДР.) И С ВНЕШНИМИ ОБЪЕКТАМИ УПРАВЛЕНИЯ И КОНТРОЛЯ	

Оригинальное описание «Вектора» к выставке 1987 года на ВДНХ.

Особенности компьютера:

- 64КБ оперативной памяти (1/10 от той, что «хватит всем»);
- Тактовая частота процессора повышена до 3 МГц, однако некоторые команды выполнялись существенно медленнее, чем могли бы;
- Трехканальный звук на базе микросхемы КР580ВМ53 (микросхема-таймер, которая была применена в компьютере по причинам дешевизны и отсутствия доступных промышленных аналогов);

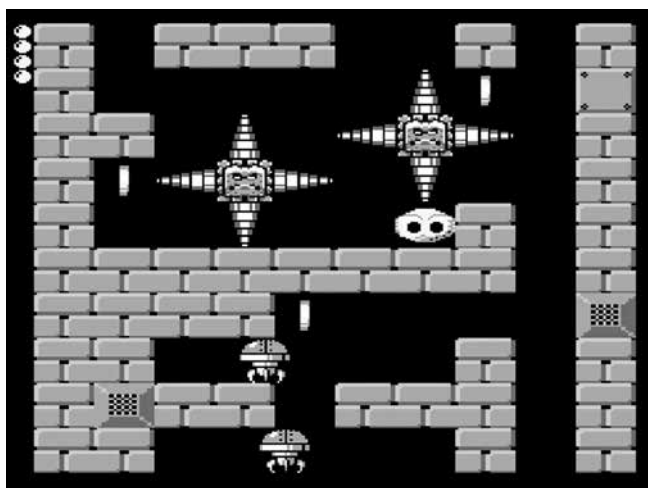
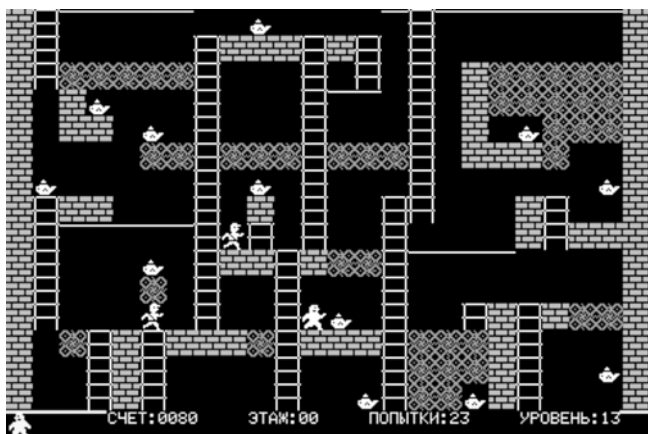
- Ввод данных через магнитофонный вход (также существовали устройства-комбодевайсы, позволяющие подключать HDD, FDD, квази-диск и AY-3-8910 для проигрывания музыки);

- Аппаратный вертикальный скроллинг;

- Три режима видео с поддержкой до 16 одновременно отображаемых цветов из палитры 256;

- Отсутствие как такового аппаратного «текстового» режима.

Для наглядности приведем несколько скриншотов с «Вектора-06Ц».



Для любителей ZX Spectrum-совместимых компьютеров кое-что в «Векторе» может стать неожиданностью, а именно отсутствие «Бейсика» в штатной прошивке. Компьютер «стартовал» сразу же готовым к загрузке с кассеты.

На сегодняшний день «Вектор-06Ц» – редкий и весьма дорогой гость в коллекциях любителей ретро. Причина тому – высокое содержание драгметаллов в компьютере. Это и позолоченные разъемы, и «дорогие» конденсаторы. Подавляющее большинство «Векторов», которые можно найти на онлайн-аукционах, либо перепаяны на менее «дефицитные» компоненты, либо пали жертвой варваров с бокорезами.

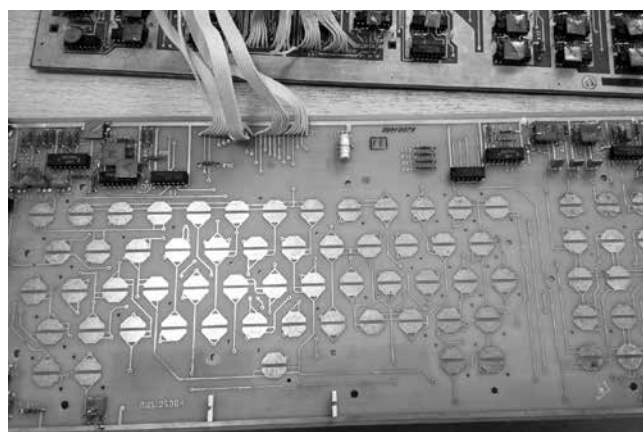
А вот с клавиатурой «Вектора» были сложности. Выпустились 2 варианта клавиатур, на герконалах и «емкостные».

Герконовые варианты были на порядок удобнее своих емкостных конкурентов (на самом деле, нажатие на

клавишу переключало две половинки «пятак»,) однако сейчас почти не встречаются. А поролон, который в комплекте с пружиной обеспечивал упругое нажатие на емкостных клавиатурах, имел тенденцию рассыпаться, в связи с чем его приходилось регулярно менять.



Герконовая клавиатура



Плата емкостной клавиатуры

Компьютер производился с 1987 по начало 1990-х годов. Постепенно был вытеснен IBM-совместимыми машинами и (локально) клонами ZX-Spectrum.

Всем, кто не знаком с этим замечательным компьютером, настоятельно рекомендуем посмотреть обзоры на уже упомянутом выше канале.

Разумеется, только «Вектором» и «ПК» применение процессора BM80A в нашей стране не ограничивается.

Другими интересными образчиками машин на базе BM80A являются шахматный компьютер «Интеллект-2», печатная машинка «ПЭЛК-3110 Элема», музыкальные синтезаторы (например, «Форманта»), игровые автоматы («ТИА-МЦ-1»), различные периферийные устройства (принтеры, УВВПЧ) и даже телефоны с АОН. Но что касается последних – здесь правил бал совершенно другой процессор, о котором речь пойдет далее.

T34BM1, KP1858BM1, KP1858BM3 и прочие аналоги Zilog Z80

Этот раздел, в общем-то, можно было бы закончить этой лаконичной картинкой...

...но это было бы слишком просто.

Оригинальный процессор Zilog Z80, родственник Intel 8080, появился на рынке в июле 1976 года. Курьезный

факт — из-за того, что Zilog свободно продавала лицензии на производство совместимых процессоров (а страны Восточной Европы и СССР игнорировали лицензирование как рудимент капитализма), Zilog в итоге выпустила менее половины от всех произведенных Z80.



Оригинальный Zilog



Отечественный клон процессора

Несмотря на то, что в СССР существовали собственные аналоги этого популярного процессора, множество популярных в то время компьютерных устройств использовали именно «оригинальные» чипы преимущественно филиппинского происхождения.

Краткие характеристики процессора:

- тактовая частота до 20 МГц;
- набор инструкций на основе i8080;
- встроенный системный контроллер;
- 8-битная шина данных;
- 16-битная шина адреса (64 Кб);
- инструкции деления и умножения отсутствуют, в ряде случаев для работы с числами применялись отдельные сопроцессоры.

Существовало несколько вариантов исполнения чипа: DIP40 и 44-контактные PLCC и PQFP.

В СССР выпуском клонов Z80 (T34BM1) занимался зеленоградский завод «Ангстрем». Чуть позже к производству новых ревизий чипа (KP1858BM1, KP1858BM3) подключились и другие заводы: воронежский «Электроника» и минский «Транзистор». Пробные экземпляры процессора сошли с конвейера в 1991 году.

В целом же это был прекрасный, мощный и недорогой процессор, на базе которого было построено немало игровых консолей, таких как Game Boy, Sega Genesis (в качестве звукового сопроцессора), компьютеров (MSX,

Amstrad CPC и пр.). А в Commodore 128 (последователей Commodore 64) Zilog Z80 использовался в качестве дополнительного ЦП для поддержки операционной системы CP/M. А если говорить о «не-компьютерном» применении процессора, список растянется не на одну страницу.

Разумеется, самым популярным вариантом применения этого процессора в нашей стране было так называемое «спектрумостроение», т.е. разработка ZX Spectrum-совместимых машин разного уровня качества и навороченности. Помимо этого, активно продавались стационарные домашние телефоны с автоматическим определителем номера (АОН) на базе Z80.



Типовой телефон с АОН на базе клона Z80

Неизвестно, является ли это отечественной разработкой, но в начале 1990-х на базе Zilog Z80 и AY-3-8910 делались музыкальные чиптюн-звонки. Подобные проекты-конструкторы для самостоятельной сборки можно найти и сейчас.

Примечательно, что в оригинальном ZX Spectrum существенная часть видеотракта и некоторые дополнительные логические модули были реализованы в проприетарной микросхеме ULA (Uncommitted Logic Array). Отечественный «ответ» ULA, микросхема T34BG1 (и ее родственники KA1515XM1, KB01BG1-2 и И185) выпускались в начале-середине 1990-х и были призваны максимально облегчить сборку клонов «Спектрума».



T34BG1 в естественной среде обитания

Кроме того, достоверно известно, что в некоторых школах, оборудованных компьютерными классами, устанавливались компьютеры MSX японского (Yamaha) производства под маркировкой КУВТ-2.



Один из учебных компьютеров MSX

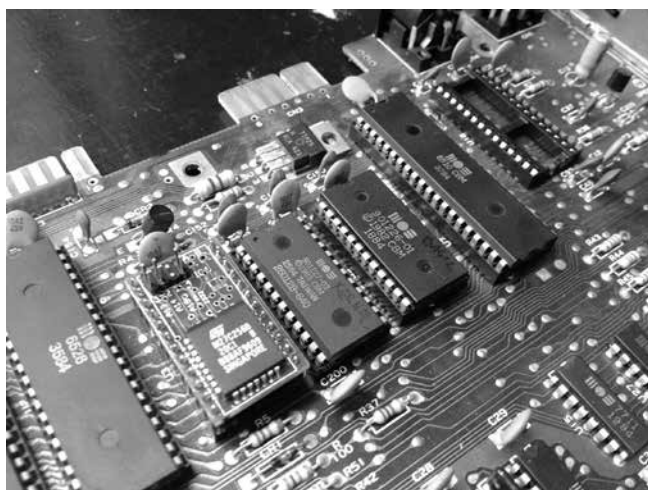
В заключение раздела приведем несколько интересных ссылок:

- Канал sinc LAIR на YouTube: крупнейший русскоязычный канал, посвященный ZX Spectrum
- Видео, посвященное сборке клона ZX Spectrum
- Демо, написанное для русского телефона с АОН на базе Z80
- 6502 и все-все-все

В 1980-х миром правили уже упомянутый Zilog Z80 и его не менее успешный и популярный «коллега» MOS Technology 6502.



Оригинальный CPU собственной персоной



Материнская плата Commodore 64

Этот восьмибитный процессор, представленный в 1975 году, имел множество ответвлений и продолжателей. Специальные версии чипа устанавливались в компьютерах от Commodore, Atari и домашних приставках Famicom/NES. Одним из немаловажных факторов успеха ЦП была его дешевизна. В то время, как Intel 8080 продавались за \$179, 6502 стоил всего \$25.

С точки зрения технических характеристик MOS 6502 крайне похож на Z80, а его архитектура в целом близка к RISC.

В нашу страну этот процессор (разумеется, в виде клона SM630) добрался на борту болгарского компьютера «Правец», клона Apple II.



«Правец»

Кроме того, производившийся с 1984 по 1990 гг. учебный компьютер «Агат», частично совместимый с «Правцем», также нес на борту процессор SM630P, еще один клон MOS6502. Поговорим об этих машинах подробнее.



Одна из версий «Агата»

8-битные версии компьютера «Правец»

Фактически, за исключением некоторых моделей, восьмибитные «Правцы» были клонами популярных американских компьютеров Apple II (Plus, e, c). Разработчик первой версии компьютера – Иван Марангозов. В годы актуальности «Правцев» коллеги шутили, что аббревиа-

тура ИМКО (Индивидуальный МикроКОмпьютер), которой маркировались некоторые модели, расшифровывается не иначе как «Иван Марангозов Копирует Оригинал».



Apple II



«Правец»

Линейка производилась с 1984 по 1994 годы. Одна из версий, «Правец 8М», известный также под именем «ИМКО-2М», прямых аналогов среди Apple-машин не имел, зато был совместим с популярной ОС CP/M и имел особый, цельно-корпусный вариант исполнения для военной промышленности. Представлен компьютер был в том же 1984-ом на Международном симпозиуме по робототехнике в Лондоне.

Что касается «модификаций» – ПЗУ компьютера было слегка модифицировано для поддержки кириллических шрифтов. При запуске на экране загоралась надпись, содержащая название компьютера.

Поздние версии компьютера собирались, помимо исторической родины, еще и в Ташкенте (некоторые модели — на Тайване) и поставлялись в учебные классы.

Популярность в школах была обусловлена тем, что при использовании специальных карт U-LAN становилось возможно объединить сразу несколько компьютеров в локальную сеть.



Архивная фотография, на которой запечатлен человек, играющий с «Правцем» в шахматы

В качестве операционной системы использовались Apple DOS, ProDOS и CP/M (для последней требовалось установить дополнительную карту).

Отдельного упоминания достоин «Правец 8D», являющийся клоном... нет, не Apple II, а Oric Atmos, самобытного британского «убийцы ZX Spectrum». Однако широкого распространения этот компьютер, как, кстати, и его «идейный вдохновитель», не получил.



«Агат»

Очень часто исследователи истории «Агата» ссылаются на статью журнала BYTE от 1984 года. В ней глазной хирург и разработчик программного обеспечения Лео Борс пишет о своих впечатлениях от работы на «плохом советском клоне Apple II». Так, поработав с компьютером пару-тройку минут, он окрестил его словом «yablochka». И не мудрено: тестовый образец «Агата» был массивен, выкрашен в «революционный красный цвет» и непрерывно кряхтел дисководом.

«Агат» можно назвать первым советским серийным компьютером. Он был разработан в НИИ Вычислительных Комплексов под руководством небезызвестного А. Ф. Иоффе в 1981-1983 гг. Серийное производство компьютеров стартовало в 1984 году.



Стильный, brutальный «Агат»



Более распространенная версия компьютера



А. Иoffee, автор «Агата»

Несмотря на неплохую совместимость с Apple II, компьютер нельзя однозначно назвать его полным клоном. Как минимум, потому что использование иностранной элементной базы без веских причин в те годы было недопустимо и конструкторы, скрепя сердце, вынуждены были создать принципиально новую плату на основе процессора серии 588, серийно производимом на минском заводе «Интеграл». Совместимость с командами 6502 достигалась сугубо средствами эмуляции. Из-за этого компьютер физически не мог тягаться по скорости с оригиналом. Именно эту версию имел удовольствие препарировать доктор Борс во время своего краткого визита в СССР.

Серийные образцы компьютера выпускались уже на оригинальном 6502 – инженерам удалось убедить руководство в целесообразности использования «чужеродного элемента». Тем не менее, в целом архитектура «Агата» во много страдает от технических ограничений отечественной элементной базы. К ним относятся, например,

использование дополнительных плат с памятью (из-за дефицита на РУ5 применялись микросхемы меньшего объема) и схема знакогенератора (только заглавные буквы). Серийные «Агаты» продавались приблизительно за 3900 рублей и были доступны как рядовым пользователям, так и средним учебным заведениям (в формате КУВТ). Немалым преимуществом «Агата» было и то, что он с завода комплектовался пятидюймовым дисководом.

Чуть позднее на ниве оборудования компьютерных классов «Агат» потеснили «Корветы» на базе все того же КР580ВМ80А, «БК-0010» и «ДВК-2».

Претерпев немало модификаций и улучшений, компьютер официально отправился на покой в 1993 году. По некоторым данным, последнюю машину выпустил Загорский электромеханический завод (ЗЭМЗ).

habr.com

HORN
TRADE
ООО «ГорнТрейд»

поставка электронных компонентов

контрактное производство

+375 17 317-92-95

+375 17 317-92-98

e-mail: info@horntrade.net

УНП 190491237

ПЯТЬ ПРЕДСКАЗАНИЙ IBM О ТОМ, КАКОЕ БУДУЩЕЕ НАС ЖДЕТ В 2022 ГОДУ

Технологический гигант IBM уже далеко не первый год делает громкие прогнозы, касающиеся «нашего» технологического будущего. И следует отметить, что процент верных предсказаний удивительным образом всегда оказывается выше, чем процент неверных. Сейчас мы рассмотрим предсказания, сделанные компанией пять лет назад, в 2017 году. И посмотрим, насколько точно они были реализованы.

■ НИКОЛАЙ ХИЖНЯК

Согласно этим предсказаниям, в течение всего нескольких лет нас ожидает серьезный рост развития искусственного интеллекта (ИИ), сверхмощных телескопов, умных сенсоров и медицинских устройств. При этом выгода от этих инноваций будет заметна в очень многих сферах, начиная от системы здравоохранения и охраны окружающей среды и заканчивая сферами, в которых разбираются вопросы нашего понимания Земли и всей окружающей Вселенной в целом.

Разумеется, следует сразу уточнить, что представленные ниже предсказания основываются лишь на тех технологиях и инновациях, к которым мы уже пришли. Поэтому следует понимать, что за взятый в этих предсказаниях временной промежуток могут случиться и другие открытия.

Может ли ИИ заменить психолога

Уже сейчас вы можете рассказать многое о человеке, просто наблюдая за тем, как он разговаривает. В голосе четко ощущается усталость, возбуждение, растерянность, невниманье или печаль.

Люди в процессе своей эволюции сами научились понимать признаки, указывающие на эти особенности состояния, однако нынешний бум развития компьютерных вычислений может, помимо прочего, привести к тому, что технологии, применяемые в анализе речи, могут стать еще более совершенными и, если можно так выразиться, более проницательными, чем мы.

«То, как мы говорим и как пишем, будет использоваться в качестве индикаторов нашего психологического состояния и физического здоровья», — делает предсказание IBM.

Как это будет происходить? Например, психологические расстройства и заболевания, такие как болезнь Паркинсона, можно будет определять с помощью обычного мобильного приложения для вашего смартфона, который будет синхронизироваться с ИИ-системой в облаке. Благодаря более раннему обнаружению признаков заболеваний, у нас будет больше времени для того, чтобы разработать средства для их лечения.

На первый взгляд связь между особенностями речи и симптомами заболеваний может показаться слишком натянутой, однако экспериментальные системы, способ-

ные выполнять эту функцию, начинают появляться уже сейчас. Например, в 2016 году команда ученых из Университета Южной Калифорнии создала программу, которая способна определять вариативность нормальной речи и сравнивать ее с признаками депрессивного состояния.

Может ли ИИ изменить зрение человека

Наши глаза сами по себе являются прекрасным примером удивительных возможностей природы и биологии, однако, согласно IBM, благодаря мощным, но при этом крошечным камерам, работающим в тандеме со сверхбыстрыми системами на базе ИИ-алгоритмов, люди обретут возможность в буквальном смысле видеть больше, чем доступно обычно человеческому глазу.

Наряду с видимым светом, люди начнут видеть микроволны, волны миллиметрового диапазона и даже изображения, получаемые в инфракрасном диапазоне спектра. И все благодаря достаточно компактным устройствам, способным поместиться в карманах обычной одежды. По сути, мы получим персональные интроскопы (устройства, применяемые для просвета багажа в тех же аэропортах), но размером с ваш телефон.

Используя подобные устройства, мы сможем безошибочно определять, например, является ли та или иная пища безопасной для нас; самоуправляемые автомобили, оснащенные такими системами, станут гораздо эффективнее «смотреть сквозь» туман или сильный дождь. И это лишь самая малая часть примеров.

Могут ли новые технологии сделать космические путешествия реальностью?

И что удивительно, одни из первых таких устройств уже появятся. Взять, к примеру, очки EnChroma, раскрашивающие не только в правильные цвета, а вообще в цвета окружающий мир для людей, страдающих дальтонизмом и ахроматопсией. Для последних мир в буквальном смысле представлен черно-белым изображением. На момент написания прогноза такие технологии очень дорогие и экспериментальные, но в 2022 году, по мнению IBM, их сможет позволить себе абсолютно любой человек.

Что такое «макроскоп» и как он поможет изучить Землю

Спутниковыми технологиями, благодаря которым мы способны заглянуть в скважину дверного замка, сейчас

уже мало кого удивишь. Однако такие сервисы, как Google Earth, – это только начало.

IBM предсказывает, что «макроскопные системы» — работающие по принципу микроскопов, только в обратную сторону — позволят объединить вместе «все комплексные данные о Земле» и наделить нас возможностью взглянуть на всю эту информацию совершенно с других точек зрения.

Эта технология не только расширит возможности сбора информации с помощью спутников, но и обеспечит нас более «умными» сенсорами, погодными станциями, а также позволит подобрать гораздо более оптимальные и структурированные способы обработки информации.

Эта технология принесет пользу не только в вопросах изучения природных процессов, происходящих на Земле и за ее пределами, но и в нашей обычной жизни. Абсолютно все виды устройств, включая наши системы освещения и даже холодильники, могут быть изучены с помощью макроскопных систем будущего, чтобы в конечном итоге мы могли комбинировать и применить эти знания в предсказаниях и предвидениях всего — начиная от климатических изменений и заканчивая поиском решений на вопросы распределения еды на планете.

От удаленно управляемых ламп к умным устройствам, которые могут записывать читаемый вами список необходимых покупок, самостоятельно заказывать для вас еду, заходить в «Википедию» и искать за вас нужную информацию... В общем, в мире, где все больше устройств, домов и даже городов становится частью общей умной инфраструктуры, подключенной к Сети, — только представьте, какой потенциал откроет обладание технологией, которая сможет собирать в буквальном смысле абсолютно всю информацию об этом мире и направлять ее на решение конкретных задач.

Какие технологии приведут к революции в медицине

С развитием и увеличением производительности компьютерных технологий, медицина является одним из тех направлений, которое получит наибольшую пользу от этого прогресса, считают в IBM. Только представьте, что вы сможете недорого и самостоятельно ставить точные медицинские диагнозы у себя дома или предсказывать приближение болезни гораздо раньше, чем сейчас.

«Новые «лаборатории на чипе» (они же микросистемы полного анализа) будут играть роль нанотехнологических докторов-детективов, отслеживающих имеющиеся в наших организмах и жидкостях следы грядущих заболеваний и давая нам немедленно знать о том, что пора посетить своего врача», — объясняют в IBM.

Другими словами, полноценные биохимические лаборатории будущего смогут уместиться на ладонях ваших рук.

Определение на ранних стадиях таких болезней, как рак или болезнь Паркинсона, может сыграть большую разницу в успешности лечения. Именно поэтому ученые работают над тем, как усовершенствовать и одновременно упростить методы проведения анализа продуктов наших выделительных систем: наших слез, крови, мочи и пота.

Наши системы мониторинга сна и фитнес-браслеты смогут отправлять данные в облачные ИИ-сервера, где будет проводиться ее эффективная и быстрая обработка. Затем эта же информация будет возвращаться к нам обратно в виде детальных советов о том, как можно улучшить наше самочувствие и здоровье. И все это будет сопровождаться одновременным автоматическим оповещением вашего лечащего врача при любых признаках приближающейся болезни.

Как быстрее обнаружить загрязнение окружающей среды с помощью ИИ

Комбинация умного оборудования и систем ИИ-анализа смогут также давать более точные и своевременные предсказания в вопросах угроз загрязнений окружающей среды, считают в IBM. Как и смарт-трекеры, которые однажды смогут указывать на ранние признаки заболеваний, смарт-сенсоры, помещенные, например, в землю или установленные на летающие дроны, смогут определять выбросы и загрязнения в реальном времени без необходимости предварительного проведения анализа образцов и данных в лабораториях.

Одним из примеров подобных загрязнений является повышение уровня метана в атмосфере — невидимое невооруженным глазом, но при этом рассматриваемое второй крупнейшей причиной глобального потепления, после углекислого газа. Смарт-сенсоры, установленные вдоль газопроводных труб, возле складских хозяйств и естественных источников этих выбросов смогут практически мгновенно определять и оповещать о повышении концентрации этого опасного газа в атмосфере.

«Такие утечки можно будет определять в течение нескольких минут, а не недель, как это происходит сейчас. Это не только позволит сократить объем утечек, но также и избежать потенциальных катастрофических последствий», — объясняет IBM.

Как уже не раз говорилось, предсказание будущего — не всегда дело благородное, да и не очень простая задача. Однако все описанные сегодня технологии уже так или иначе находятся в разработке исследовательских и научных команд по всему миру. Другими словами, то, когда мы получим эти технологии, — остается лишь вопросом времени.

ibm.com

ИСКУССТВЕННЫЙ ИНТЕЛЛЕКТ. ВСЕ, ЧТО ВЫ ХОТЕЛИ ЗНАТЬ

Причина, по которой эта (и другие) статья появилась на свет, проста: возможно, искусственный интеллект – не просто важная тема для обсуждения, а самая важная в контексте будущего. Все, кто хоть немного проникает в суть потенциала искусственного интеллекта, признают, что оставлять без внимания эту тему нельзя. Некоторые – и среди них Элон Маск, Стивен Хокинг, Билл Гейтс, не самые глупые люди нашей планеты – считают, что искусственный интеллект представляет экзистенциальную угрозу для человечества, сопоставимую по масштабам с полным вымиранием нас как вида. Что ж, усаживайтесь поудобнее и расставляйте для себя все точки над i.

■ **ИЛЬЯ ХЕЛЬ**

ЧАСТЬ ПЕРВАЯ: ПУТЬ К СВЕРХИНТЕЛЛЕКТУ

Будущее ИИ

Представьте, что машина времени перенесла вас в 1750 год – во времена, когда мир испытывал постоянные перебои с поставками электричества, связь между городами подразумевала выстрелы из пушки, а весь транспорт работал на сене. Допустим, вы попадаете туда, забираете кого-нибудь и приводите в 2015 год, показать, как оно тут все. Мы не в состоянии понять, каково ему было бы увидеть все эти блестящие капсулы, летящие по дорогам; поговорить с людьми по ту сторону океана; посмотреть на спортивные игры за тысячу километров от него; услышать музыкальное выступление, записанное 50 лет назад; поиграть с волшебным прямоугольником, который может сделать снимок или запечатлеть живой момент; построить карту с паранормальной голубой точкой, обозначающей его местоположение; смотреть на чье-то лицо и общаться с ним за много километров и так далее. Все это – необъяснимое волшебство для почти трехсотлетних людей. Не говоря уже об Интернете, Международной космической станции, Большом адронном коллайдере, ядерном оружии и общей теории относительности.

Такой опыт для него не будет удивительным или шокирующим – эти слова не передают всей сути мысленного коллапса. Наш путешественник вообще может умереть.

Но есть интересный момент. Если он вернется в 1750 год и ему станет завидно, что мы захотели взглянуть на его реакцию на 2015 год, он может захватить с собой машину времени и попробовать проделать то же самое, скажем, с 1500 годом. Прилетит туда, найдет человека, заберет в 1750 год и все покажет. Парень из 1500 года будет шокирован безмерно – но вряд ли умрет. Хотя он, конечно, будет удивлен, разница между 1500 и 1750 годом куда меньше, чем между 1750 и 2015. Человек из 1500 года удивится некоторым моментам из физики, поразится тому, какой стала Европа под жесткой пятой империализма, нарисует в голове новую карту мира. Но повседневная жизнь 1750 года – транспорт, связь и т. п. – вряд ли удивит его до смерти.

Нет, чтобы парень из 1750 года повеселился так же, как мы с ним, он должен отправиться куда дальше – возможно, год так в 12 000 до н. э., еще до того, как первая сельскохозяйственная революция позволила зародиться

первым городам и понятию цивилизации. Если бы кто-либо из мира охотников-собирателей, со времен, когда люди в большей степени были еще очередным животным видом, увидел огромные человеческие империи 1750 года с их высокими церквями, судами, пересекающими океаны, их понятие быть «внутри» здания, все эти знания – он бы умер, скорее всего.

И тогда, после смерти, он бы позавидовал и захотел сделать то же самое. Вернулся бы на 12 000 лет назад, в 24 000 год до н. э., забрал бы человека и притащил его в свое время. И новый путешественник сказал бы ему: «Ну такое, хорошо, спасибо». Потому что в этом случае человеку из 12 000 год до н. э. нужно было бы вернуться на 100 000 лет назад и показать местным аборигенам огонь и язык в первый раз.

Если нам нужно перевезти кого-то в будущее, чтобы тот был до смерти удивлен, прогресс должен пройти определенное расстояние. Должна быть достигнута Точка Смертельного Прогресса (ТСП). То есть, если во времена охотников-собирателей ТСП занимала 100 000 лет, следующая остановка состоялась уже в 12 000 году до н. э. За ней прогресс шел уже быстрее и кардинально преобразил мир к 1750 году (ориентировочно). Потом понадобилось пару сотен лет, и вот мы здесь.

Эту картину – когда человеческий прогресс движется быстрее по мере течения времени – футуролог Рэй Курцвейл называет законом ускоряющейся отдачи человеческой истории. Это происходит, потому что у более развитых обществ есть возможность двигать прогресс более быстрыми темпами, нежели у менее развитых обществ. Люди 19 века знали больше, чем люди 15 века, поэтому неудивительно, что прогресс в 19 веке шел более стремительными темпами, нежели в 15 веке, и так далее.

На меньших масштабах это тоже работает. Фильм «Назад в будущее» вышел в 1985 году, и «прошлое» было в 1955 году. В фильме, когда Майкл Джей Фокс вернулся в 1955 год, его застали врасплох новизна телевизоров, цены на содовую, отсутствие любви к гитарному звуку и вариации в сленге. Это был другой мир, безусловно, но если бы фильм снимали сегодня, а прошлое было в 1985 году, разница была бы куда более глобальна. Марти Макфлай, попавший в прошлое из времени персональных компьютеров, Интернета, мобильных телефонов, был бы



«Мы стоим на пороге изменений, сравнимых с зарождением человеческой жизни на Земле» (Вернор Виндж).

гораздо более неуместным, чем Марти, отправившийся в 1955 год из 1985.

Все это благодаря закону ускоряющейся отдачи. Средняя скорость развития прогресса между 1985 и 2015 годами была выше, чем скорость с 1955 по 1985 годы – потому что в первом случае мир был более развитым, он был насыщен достижениями последних 30 лет.

Таким образом, чем больше достижений, тем быстрее происходят изменения. Но разве это не должно оставлять нам определенные намеки на будущее?

Курцвейл предполагает, что прогресс всего 20 века мог бы быть пройден всего за 20 лет при уровне развития 2000 года – то есть в 2000 году скорость прогресса была в пять раз выше, чем средняя скорость прогресса 20 века. Также он считает, что прогресс всего 20 века был эквивалентен прогрессу периода с 2000 по 2014 год, и прогресс еще одного 20 века будет эквивалентен периоду до 2021 года – то есть всего за семь лет. Спустя несколько десятков лет весь прогресс 20 века будет проходиться по несколько раз в год, а дальше – всего за месяц. В конечном итоге закон ускоряющейся отдачи доведет нас до того, что за весь 21 век прогресс будет в 1000 раз превышать прогресс 20 века.

Если Курцвейл и его сторонники правы, 2030 год удивит нас так же, как парня из 1750 года удивил бы наш 2015 – то есть следующая ТСП займет всего пару десятков лет – а мир 2050 года будет так сильно отличаться от современного, что мы едва ли его узнаем. И это не фантастика. Так полагает множество ученых, которые умнее и образованнее нас с вами. И если вы взглянете

в историю, то поймете, что это предсказание вытекает из чистой логики.

Почему же, когда мы сталкиваемся с заявлениями вроде «мир через 35 лет изменится до неузнаваемости», мы скептически пожимаем плечами? Есть три причины нашего скепсиса относительно прогнозов будущего:

1. Когда дело доходит до истории, мы думаем прямыми цепочками. Пытаясь представить прогресс следующих 30 лет, мы смотрим на прогресс предыдущих 30 как на индикатор того, сколько всего, по всей вероятности, произойдет. Когда мы думаем о том, как изменится наш мир в 21 веке, мы берем прогресс 20 века и добавляем его к 2000 году. Такую же ошибку совершает наш парень из 1750 года, когда достает кого-то из 1500 года и пытается его удивить. Мы интуитивно думаем линейным образом, хотя должны бы экспоненциальным. По существу, футуролог должен пытаться предугадать прогресс следующих 30 лет, не глядя на предыдущие 30, а судя по текущему уровню прогресса. Тогда прогноз будет точнее, но все равно мимо ворот.

Чтобы думать о будущем корректно, вам нужно видеть движение вещей в куда более быстром темпе, чем они движутся сейчас.

2. Траектория недавней истории зачастую выглядит искаженно. Во-первых, даже крутая экспоненциальная кривая кажется линейной, когда вы видите небольшие ее части. Во-вторых, экспоненциальный рост не всегда гладкий и однородный. Курцвейл считает, что прогресс движется змееподобными кривыми.

Такая кривая проходит через три фазы: 1) медленный рост (ранняя фаза экспоненциального роста); 2) быстрый

рост (взрывная, поздняя фаза экспоненциального роста); 3) стабилизация в виде конкретной парадигмы.

Если вы взглянете на последнюю историю, часть S-кривой, в которой вы в данный момент находитесь, может скрывать от вашего восприятия скорость прогресса. Часть времени между 1995 и 2007 годами ушла на взрывное развитие Интернета, представление Microsoft, Google и Facebook публике, рождению социальных сетей и развитию сотовых телефонов, а затем и смартфонов. Это была вторая фаза нашей кривой. Но период с 2008 по 2015 год был менее прорывным, по крайней мере на технологическом фронте. Те, кто думают о будущем сегодня, могут взять последние пару лет для оценки общего темпа прогресса, но они не видят большей картины. На деле же новая и мощная фаза 2 может назреть уже сейчас.

3. Наш собственный опыт делает нас ворчливыми стариками, когда речь заходит о будущем. Мы основываем свои идеи о мире на собственном опыте, и этот опыт установил темпы роста в недавнем прошлом для нас как «само собой разумеющееся». Также и наше воображение ограничено, поскольку использует наш опыт для прогнозирования – но чаще всего у нас просто нет инструментов, которые позволяют точно предположить будущее. Когда мы слышим прогнозы на будущее, которые расходятся с нашим повседневным восприятием работы вещей, мы инстинктивно считаем их наивными. Если бы я сказал вам, что вы доживете до 150 или 250 лет, а может быть, и вовсе не умрете, вы инстинктивно подумаете, что «это глупо, я знаю из истории, что за это время умирали все». Так и есть: никто не доживал до таких лет. Но и ни один самолет не летал до изобретения самолетов.

Таким образом, хотя скепсис кажется вам обоснованным, чаще всего он ошибочен. Нам стоит принять, что если мы вооружаемся чистой логикой и ждем привычных исторических зигзагов, мы должны признать, что очень, очень и очень многое должно измениться в ближайшие десятилетия; намного больше, чем можно предположить интуитивно. Логика также подсказывает, что если самый продвинутый вид на планете продолжает делать гигантские скачки вперед, быстрее и быстрее, в определенный момент скачок будет настолько серьезным, что он кардинально изменит жизнь, какой мы ее знаем. Нечто подобное случилось в процессе эволюции, когда человек стал настолько умен, что полностью изменил жизнь любого другого вида на планете Земля. И если вы потратите немного времени на чтение того, что происходит сейчас в науке и технике, вы, возможно, начнете видеть определенные подсказки о том, каким будет следующий гигантский скачок.

Что такое ИИ (искусственный интеллект)?

Как и многие на этой планете, вы привыкли считать искусственный интеллект глупой идеей научной фантастики. Но за последнее время очень много серьезных людей проявило обеспокоенность этой глупой идеей. Что не так?

Есть три причины, которые приводят к путанице вокруг термина ИИ:

- Мы ассоциируем ИИ с фильмами. «Звездные войны». «Терминатор». «Космическая Одиссея 2001 года». Но как и роботы, ИИ в этих фильмах – вымысел. Таким образом, голливудские ленты разбавляют уровень нашего восприятия, ИИ становится привычным, родным и, безусловно, злым.

- Это широкое поле для применения. Оно начинается с калькулятора в вашем телефоне и разработки самоуправляемых автомобилей и доходит до чего-то далекого в будущем, что кардинально изменит мир. ИИ обозначает все эти вещи, и это сбивает с толку.

- Мы используем ИИ каждый день, но зачастую даже не отдаем себе в этом отчета. Как говорил Джон Маккарти, изобретатель термина «искусственный интеллект» в 1956 году, «как только он заработал, никто больше не называет его ИИ». ИИ стал больше как мифическое предсказание о будущем, нежели что-то реальное. В то же время есть в этом названии и привкус чего-то из прошлого, что никогда не стало реальностью. Рэй Курцвейл говорит, что он слышит, как люди ассоциируют ИИ с фактами из 80-х годов, что можно сравнить с «утверждением, что интернет умер вместе с доткомом в начале 2000-х».

Давайте проясним. Во-первых, перестаньте думать о роботах. Робот, который является контейнером для ИИ, иногда имитирует человеческую форму, иногда нет, но сам ИИ – это компьютер внутри робота. ИИ – это мозг, а робот – тело, если у него вообще есть это тело. К примеру, программное обеспечение и данные Siri – это искусственный интеллект, женский голос – персонификация этого ИИ, и никаких роботов в этой системе нет.

Во-вторых, вы наверняка слышали термин «сингулярность» или «технологическая сингулярность». Этот термин используется в математике для описания необычной ситуации, когда обычные правила больше не работают. В физике он используется для описания бесконечно малой и плотной точки черной дыры или изначальной точки Большого Взрыва. Опять же, законы физики в ней не работают. В 1993 году Вернор Виндж написал знаменитое эссе, в котором применил этот термин к моменту в будущем, когда интеллект наших технологий превзойдет наш собственный – и в этот момент жизнь, какой мы ее знаем, изменится навсегда, а обычные правила ее существования больше не будут работать. Рэй Курцвейл в дальнейшем уточнил этот термин, указав, что сингулярность будет достигнута, когда закон ускоряющейся отдачи достигнет экстремальной точки, когда технологический прогресс будет двигаться так быстро, что мы перестанем замечать его достижения, почти бесконечно быстро. Тогда мы будем жить в совершенно новом мире. Однако многие эксперты перестали использовать этот термин, поэтому давайте и мы не будем часто к нему обращаться.

Наконец, хотя есть много типов или форм ИИ, которые вытекают из широкого понятия ИИ, основные категории его зависят от калибра. Есть три основных категории:

- Узконаправленный (слабый) искусственный интеллект (УИИ). УИИ специализируется в одной области. Среди таких ИИ есть те, кто может обыграть чемпиона

мира по шахматам, но на этом все. Есть такой, который может предложить лучший способ хранения данных на жестком диске, и все.

- **Общий (сильный) искусственный интеллект.** Иногда также называют ИИ человеческого уровня. ОИИ относят к компьютеру, который умен, как человек – машина, которая способна выполнять любое интеллектуальное действие, присущее человеку. Создать ОИИ намного сложнее, чем УИИ, и мы пока до этого не дошли. Профессор Линда Готтфредсон описывает интеллект как «в общем смысле психический потенциал, который, наряду с другими вещами, включает способность рассуждать, планировать, решать проблемы, мыслить абстрактно, понимать сложные идеи, быстро учиться и извлекать опыт». ОИИ должен уметь делать все это так же просто, как делаете вы.

- **Искусственный сверхинтеллект (ИСИ).** Оксфордский философ и теоретик ИИ Ник Бостром определяет сверхинтеллект как «интеллект, который гораздо умнее лучших человеческих умов в практически любой сфере, включая научное творчество, общую мудрость и социальные навыки». Искусственный сверхинтеллект включает в себя как компьютер, который немного умнее человека, так и тот, который в триллионы раз умнее в любом направлении. ИСИ и есть причина того, что растет интерес к ИИ, а также того, что в таких обсуждениях часто появляются слова «вымирание» и «бессмертие».

В настоящее время люди уже покорили самую первую ступень калибра ИИ – УИИ – во многих смыслах. Революция ИИ – это путь от УИИ через ОИИ к ИСИ. Этот путь мы можем не пережить, но он, определенно, изменит все.

Давайте внимательно разберем, как видят этот путь ведущие мыслители в этой области и почему эта революция может произойти быстрее, чем вы могли бы подумать.

Человек и искусственный интеллект

Узконаправленный искусственный интеллект – это машинный интеллект, который равен или превышает человеческий интеллект или эффективность в выполнении определенной задачи.

Несколько примеров:

- Автомобили битком набиты системами УИИ, от компьютеров, которые определяют, когда должна заработать антиблокировочная тормозная система, до компьютера, который определяет параметры системы впрыска топлива. Самоуправляемые автомобили Google, которые в настоящее время проходят испытания, будут содержать надежные системы УИИ, которые будут воспринимать и реагировать на мир вокруг себя.

- Ваш телефон – маленькая фабрика УИИ. Когда вы используете приложение карт, получаете рекомендации по скачиванию приложений или музыки, проверяете погоду на завтра, говорите с Siri или делаете что-либо еще – вы используете УИИ.

- Спам-фильтр вашей электронной почты – классический тип УИИ. Он начинает с выяснения того, как отделить спам от пригодных писем, а затем обучается в процессе обработки ваших писем и предпочтений.

- А это неловкое чувство, когда еще вчера вы искали шуруповерт или новую плазму в поисковике, а сегодня видите предложения услужливых магазинов на других сайтах? Или когда в социальной сети вам рекомендуют добавить интересных людей в друзья? Все это системы УИИ, которые работают вместе, определяя ваши предпочтения, выуживая из Интернета данные о вас, подбираясь к вам все ближе и ближе. Они анализируют поведение миллионов людей и делают выводы на основе этих анализов так, чтобы продавать услуги крупных компаний или делать их сервисы лучше.

- Google Translate – еще одна классическая система УИИ, впечатляюще хороша в определенных вещах. Распознавание голоса – тоже. Когда ваш самолет садится, терминал для него определяет не человек. Цену билета – тоже. Лучшие в мире шашки, шахматы, нарды, балды и прочие игры сегодня представлены узконаправленными искусственными интеллектами.

- Поиск Google – это один гигантский УИИ, который использует невероятно хитроумные методы для ранжирования страниц и определения результатов поисковой выдачи.

И это только в потребительском мире. Сложные системы УИИ широко используются в военных, производственных и финансовых отраслях; в медицинских системах (вспомните Watson от IBM) и так далее.

Системы УИИ в настоящем виде не представляют угрозы. В худшем случае глючный или плохо запрограммированный УИИ может привести к локальному бедствию, создать перебои в энергоснабжении, обвалить финансовые рынки и тому подобное. Но хотя УИИ не обладает полномочиями для создания экзистенциальной угрозы, мы должны видеть вещи шире – нас ждет сокрушительный ураган, предвестником которого выступают УИИ. Каждая новая инновация в сфере УИИ добавляет один блок к дорожке, ведущей к ОИИ и ИСИ. Или как хорошо заметил Аарон Саенц, УИИ нашего мира похожи на «аминокислоты первичного бульона юной Земли» – пока неживые компоненты жизни, которые однажды проснутся.

Путь от УИИ к ОИИ: почему так сложно?

Ничто так не раскрывает сложность человеческого интеллекта, как попытка создать компьютер, который будет так же умен. Строительство небоскребов, полеты в космос, секреты Большого Взрыва – все это ерунда по сравнению с тем, чтобы повторить наш собственный мозг или хотя бы просто понять его. В настоящее время мозг человека – самый сложный объект в известной Вселенной.

Возможно, вы даже не подозреваете, в чем сложность создать ОИИ (компьютер, который будет умен, как человек, в общем, а не только в одной области). Создать компьютер, который может перемножать два десятизначных числа за долю секунды – проще простого. Создать такого, который сможет взглянуть на собаку и кошку и сказать, где собака, а где кошка – невероятно сложно. Создать ИИ, который может обыграть гроссмейстера? Сделано. Теперь попробуйте заставить его прочитать абзац из книги для шестилетних детей и не только понять слова, но и их

значение. Google тратит миллиарды долларов, пытаясь это сделать. Со сложными вещами – вроде вычислений, расчета стратегий финансовых рынков, перевода языка – с этим компьютер справляется с легкостью, а с простыми вещами – зрение, движение, восприятие – нет. Как выразился Дональд Кнут, «ИИ сейчас делает практически все, что требует «мышления», но не может справиться с тем, что делают люди и животные не задумываясь».

Когда вы задумаетесь о причинах этого, вы поймете, что вещи, которые кажутся нам простейшими в исполнении, только кажутся такими, потому что были оптимизированы для нас (и животных) в ходе сотен миллионов лет эволюции. Когда вы протягиваете руку к объекту, мышцы, суставы, кости ваших плеч, локтей и кистей мгновенно выполняют длинные цепочки физических операций, синхронных с тем, что вы видите, и движут вашу руку в трех измерениях. Вам это кажется простым, потому что за эти процессы отвечает идеальное программное обеспечение вашего мозга. Этот простой трюк позволяет сделать процедуру регистрации нового аккаунта с вводом криво написанного слова (капчи) простым для вас и адом для вредоносного бота. Для нашего мозга в этом нет ничего сложного: нужно просто уметь видеть.

С другой стороны, перемножение крупных чисел или игра в шахматы – новые виды активности для биологических существ, и у нас не было достаточно времени, чтобы совершенствоваться в них (не миллионы лет), поэтому компьютеру несложно нас одолеть. Просто подумайте об этом: что бы вы предпочли, создать программу, которая может перемножать большие числа, или программу, которая узнает букву Б в миллионах ее видов написаний, в самых непредсказуемых шрифтах, от руки или палкой на снегу?

Все, что мы только что обозначили, это еще вертушка айсберга, касающаяся понимания и обработки информации. Чтобы выйти на один уровень с человеком, компьютер должен понимать разницу в тонких выражениях лица, разницу между удовольствием, печалью, удовлетворением, радостью, и почему Чацкий молодец, а Молчалин – нет.

Что же делать?

Первый шаг к созданию ИИИ: увеличение вычислительной мощности. Одна из необходимых вещей, которая должна произойти, чтобы ИИИ стал возможным, это увеличение мощности компьютерного оборудования. Если система искусственного интеллекта должна быть такой же умной, как мозг, ей нужно сравняться с мозгом по сырой вычислительной мощности.

Один из способов увеличить эту способность заключается в общем числе вычислений в секунду (OPS), которое может производить мозг, и вы можете определить это число, выяснив максимальное число OPS для каждой структуры мозга и сведя их воедино.

Рэй Курцвейл пришел к выводу, что достаточно взять профессиональную оценку OPS одной структуры и ее вес относительно веса всего мозга, а затем умножить пропорционально, чтобы получить общую оценку. Звучит

немного сомнительно, но он проделал это много раз с разными оценками разных областей и всегда приходил к одному и тому же числу: порядка 10^{16} , или 10 квадриллионов OPS.

Самый быстрый суперкомпьютер в мире, китайский «Тяньхэ-2», уже обошел это число: он способен проделывать порядка 32 квадриллиона операций в секунду. Но «Тяньхэ-2» занимает 720 квадратных метров пространства, сжигает 24 мегаватта энергии (наш мозг потребляет всего 20 ватт) и стоит 390 миллионов долларов. О коммерческом или широком применении речь не идет.

Курцвейл предполагает, что мы оцениваем состояние компьютеров по тому, как много OPS можно купить за 1000 долларов. Когда это число достигнет человеческого уровня – 10 квадриллионов OPS – ИИИ вполне может стать частью нашей жизни.

Закон Мура – исторически надежное правило, определяющее, что максимальная вычислительная мощность компьютеров умножается вдвое каждые два года – подразумевает, что развитие компьютерной техники, как и движение человека по истории, растет по экспоненте. Если соотнести это с правилом тысячи долларов Курцвейла, мы сейчас можем позволить себе 10 триллионов OPS за 1000 долларов.

Компьютеры за 1000 долларов по своим вычислительным способностям обходят мозг мыши и в тысячу раз слабее человека. Это кажется плохим показателем, пока мы не вспомним, что компьютеры были в триллион раз слабее человеческого мозга в 1985 году, в миллиард – в 1995, и в миллион – в 2005. К 2025 году мы должны получить доступный компьютер, не уступающий по вычислительной мощности нашему мозгу.

Таким образом, сырая мощь, необходимая для ИИИ, уже технически доступна. В течение 10 лет она выйдет из Китая и распространится по миру. Но одной вычислительной мощности недостаточно. И следующий вопрос: как нам обеспечить всей этой мощью интеллект человеческого уровня?

Второй шаг к созданию ИИИ: дать ему разум. Эта часть довольно сложновыполнимая. По правде говоря, никто толком не знает, как сделать машину разумной – мы до сих пор пытаемся понять, как создать разум человеческого уровня, способный отличить кошку от собаки, выделить Б, нарисованную на снегу, и проанализировать второсортный фильм. Однако есть горстка дальновидных стратегий, и в один прекрасный момент одна из них должна сработать.

Повторить мозг

Этот вариант похож на то, будто ученые сидят в одном классе с ребенком, который очень умен и хорошо отвечает на вопросы; и даже если они усердно пытаются постигать науку, они и близко не догоняют умницу-ребенка. В конечном итоге они решают: к черту, просто спишем ответы на вопросы у него. В этом есть смысл: мы не можем создать сверхсложный компьютер, так почему бы не взять за

основу один из лучших прототипов вселенной: наш мозг?

Научный мир трудится в поте лица, пытаясь выяснить, как работает наш мозг и как эволюция создала такую сложную вещь. По самым оптимистичным оценкам, получится это у них только к 2030 году. Но как только мы поймем все секреты работы мозга, его эффективности и мощности, мы сможем вдохновиться его методами в создании технологий. К примеру, одной из компьютерных архитектур, которая имитирует работу мозга, является нейронная сеть. Она начинается с сети транзисторов «нейронов», соединенных друг с другом входом и выходом, и не знает ничего – как новорожденный. Система «обучается», пытаясь выполнять задания, распознавать рукописный текст и тому подобное. Связи между транзисторами укрепляются в случае правильного ответа и ослабляются в случае неверного. Спустя много циклов вопросов и ответов система образует умные нейронные переплетения, оптимизированные для выполнения определенных задач. Мозг обучается похожим образом, но в куда более сложной манере, и пока мы продолжаем изучать его, мы открываем новые невероятные способы улучшить нейронные сети.

Еще более экстремальный плагиат включает стратегию под названием «эмуляция полного мозга». Цель: нарезать настоящий мозг тонкими пластинками, отсканировать каждую из них, затем точно восстановить трехмерную модель, используя программное обеспечение, а затем воплотить ее в мощном компьютере. Тогда у нас будет компьютер, который официально сможет делать все, что может делать мозг: ему просто нужно будет обучаться и собирать информацию. Если у инженеров получится, они смогут эмулировать настоящий мозг с такой невероятной точностью, что после загрузки на компьютер настоящая личность мозга и его память останутся нетронутыми. Если мозг принадлежал Вадиму перед тем, как он умер, компьютер проснется в роли Вадима, который теперь будет ОИИ человеческого уровня, а мы, в свою очередь, займемся превращением Вадима в невероятно разумный ИСИ, чему он наверняка обрадуется.

Насколько мы далеки от полной эмуляции мозга? По правде говоря, мы только-только эмулировали мозг миллиметрового плоского червя, который содержит 302 нейрона в общей сложности. Мозг человека содержит 100 миллиардов нейронов. Если вам попытки добраться до этого числа кажутся бесполезными, вспомните об экспоненциальных темпах роста прогресса. Следующим шагом станет эмуляция мозга муравья, затем будет мышь, а там и до человека рукой подать.

Попытаться пройти по следам эволюции

Что ж, если мы решим, что ответы умного ребенка слишком сложны, чтобы их списать, мы можем попытаться пройти по его следам обучения и подготовки к экзаменам. Что мы знаем? Построить компьютер такой же мощный, как мозг, вполне возможно – эволюция нашего собственного мозга это доказала. И если мозг слишком сложен для эмуляции, мы можем попытаться эмулировать эволюцию.

Дело в том, что даже если мы сможем эмулировать мозг, это может быть похоже на попытку построить самолет путем нелепого махания руками, повторяющего движения крыльев птиц. Чаще всего нам удастся создать хорошие машины, используя машинно-ориентированный подход, а не точное подражание биологии.

Как имитировать эволюцию, чтобы построить ОИИ? Этот метод под названием «генетические алгоритмы» должен работать примерно так: должен быть производительный процесс и его оценка, и это будет повторяться снова и снова (точно так же биологические существа «существуют» и «оцениваются» по их способности воспроизводиться). Группа компьютеров будет выполнять задачи, а самые успешные из них будут разделять свои характеристики с другими компьютерами, «выводиться». Менее успешные будут беспощадно выбрасываться на свалку истории. Через много-много итераций этот процесс естественного отбора позволит вывести лучшие компьютеры. Сложность заключается в создании и автоматизации циклов выведения и оценки, чтобы процесс эволюции шел сам по себе.

Недостатком копирования эволюции является то, что эволюции требуются миллиарды лет, чтобы что-то сделать, а нам нужно на это всего несколько десятилетий.

Но у нас есть масса преимуществ, в отличие от эволюции. Во-первых, у нее нет дара предвидения, она работает случайно – выдает бесполезные мутации, например, – а мы можем контролировать процесс в рамках поставленных задач. Во-вторых, у эволюции нет цели, в том числе и стремления к интеллекту – иногда в окружающей среде некоторый вид выигрывает не за счет интеллекта (потому что последний потребляет больше энергии). Мы же, с другой стороны, можем нацелиться на увеличение интеллекта. В-третьих, чтобы выбрать интеллект, эволюции необходимо произвести ряд сторонних улучшений – вроде перераспределения потребления энергии клетками, – мы же можем просто убрать лишнее и использовать электричество. Вне всяких сомнений, мы будем быстрее эволюции – но опять же, непонятно, сможем ли мы ее превзойти.

Предоставить компьютеры самим себе

Это последний шанс, когда ученые совсем отчаиваются и пытаются запрограммировать программу на саморазвитие. Однако этот метод может оказаться самым перспективным из всех. Идея в том, что мы создаем компьютер, у которого будет два основных навыка: исследовать ИИ и кодировать изменения в себе – что позволит ему не только больше узнавать, но и улучшать собственную архитектуру. Мы можем обучить компьютеры быть компьютерными инженерами самим себе, чтобы они саморазвивались. И их основной задачей будет выяснить, как стать умнее. Подробнее об этом мы еще поговорим.

Все это может случиться очень скоро

Стремительное развитие аппаратного обеспечения и эксперименты с программным обеспечением текут парал-

тельно, и ОИИ может появиться быстро и неожиданно по двум основным причинам:

1. Экспоненциальный рост идет интенсивно, и то, что кажется черепашными шагами, может быстро перерасти в семимильные прыжки;

2. Когда дело доходит до программного обеспечения, прогресс может казаться медленным, но затем один прорыв мгновенно меняет скорость продвижения вперед (хороший пример: во времена геоцентрического мировосприятия, людям было сложно рассчитать работу вселенной, но открытие гелиоцентризма все значительно упростило). Или, когда дело доходит до компьютера, который улучшает сам себя, все может казаться чрезвычайно медленным, но иногда всего одна поправка в системе отделяет ее от тысячекратной эффективности по сравнению с человеком или прежней версией.

Дорога от ОИИ к ИСИ

В определенный момент мы обязательно получим ОИИ – общий искусственный интеллект, компьютеры с общим человеческим уровнем интеллекта. Компьютеры и люди будут жить вместе. Или не будут.

Дело в том, что ОИИ с таким же уровнем интеллекта и вычислительной мощности, что и человек, будет по-прежнему иметь значительные преимущества перед людьми.

Оборудование

Скорость. Нейроны мозга работают с частотой 200 Гц, в то время как современные микропроцессоры (которые значительно медленнее тех, что мы получим к моменту создания ОИИ) работают с частотой 2 ГГц, или в 10 миллионов раз быстрее наших нейронов. И внутренние коммуникации мозга, которые могут двигаться со скоростью 120 м/с, существенно уступают возможности компьютеров использовать оптику и скорость света.

Размер и хранение. Размер мозга ограничен размерами наших черепов, и он не может стать больше, в противном случае внутренним коммуникациям со скоростью 120 м/с потребуется слишком много времени, чтобы проходить от одной структуры к другой. Компьютеры могут расширяться до любого физического размера, задействовать больше оборудования, наращивать оперативную память, длительную память – все это выходит за рамки наших возможностей.

Надежность и долговечность. Не только память компьютера точнее человеческой. Компьютерные транзисторы точнее биологических нейронов и менее склонны к ухудшению (да и вообще, могут быть заменены или отремонтированы). Мозги людей быстрее устают, компьютеры же могут работать без остановки, 24 часа в сутки, 7 дней в неделю.

Программное обеспечение

Возможность редактирования, модернизации, более широкий спектр возможностей. В отличие от человеческого мозга, компьютерную программу можно легко

исправить, обновить, провести эксперимент с ней. Модернизации могут также подвергаться зоны, в которых человеческие мозги слабы. Программное обеспечение человека, отвечающее за зрение, великолепно устроено, но с точки зрения инженерии его способности все же весьма ограничены – мы видим только в видимом спектре света.

Коллективная способность. Люди превосходят другие виды в плане грандиозного коллективного разума. Начиная с развития языка и формирования крупных обществ, двигаясь через изобретения письма и печати и в настоящее время активизируясь с помощью таких инструментов, как интернет, коллективный разум людей является важной причиной, по которой мы можем величать себя венцом эволюции. Но компьютеры все равно будут лучше. Всемирная сеть искусственных интеллектов, работающих на одной программе, постоянно синхронизирующихся и саморазвивающихся, позволит мгновенно добавлять в базу новую информацию, где бы ее ни достали. Такая группа также сможет работать над одной целью, как одно целое, потому что компьютеры не страдают от наличия особого мнения, мотивации и личной заинтересованности, как люди.

ИИ, который, вероятнее всего, станет ОИИ посредством запрограммированного самосовершенствования, не увидит «интеллект человеческого уровня» как важную веху – эта веха важна только для нас. У него не будет никаких причин останавливаться на этом сомнительном уровне. А учитывая те преимущества, которыми будет обладать даже ОИИ человеческого уровня, довольно очевидно, что человеческий интеллект станет для него короткой вспышкой в гонке за превосходством в интеллектуальном плане.

Такое развитие событий может нас весьма и весьма удивить. Дело в том, что, с нашей точки зрения, а) единственный критерий, который позволяет нам определять качество интеллекта, это интеллект животных, который по умолчанию ниже нашего; б) для нас самые умные люди ВСЕГДА умнее самых глупых. Примерно так:

То есть пока ИИ просто пытается достичь нашего уровня развития, мы видим, как он становится умнее, приближаясь к уровню животного.

Когда он доберется до первого человеческого уровня – Ник Бостром использует термин «деревенский дурачок», – мы будем в восторге: «Ух ты, он уже как debil. Круть!». Единственное, что заключается в том, что в общем спектре интеллекта людей, от деревенского дурачка до Эйнштейна, диапазон невелик – поэтому после того, как ИИ доберется до уровня дурачка и станет ОИИ, он внезапно станет умнее Эйнштейна.

И что будет дальше?

Взрыв интеллекта

Надеюсь, вам было интересно и весело, потому что именно с этого момента обсуждаемая нами тема становится ненормальной и жуткой. Нам стоит сделать паузу и напомнить себе, что каждый указанный выше и дальше факт – реальная наука и реальные прогнозы на будущее,

высказанные самыми выдающимися мыслителями и учеными. Просто имейте в виду.

Итак, как мы обозначили выше, все наши современные модели по достижению ОИИ включают вариант, когда ИИ самосовершенствуется. И как только он становится ОИИ, даже системы и методы, с помощью которых он вырос, становятся достаточно умны, чтобы самосовершенствоваться – если захотят. Возникает интересное понятие: рекурсивное самосовершенствование. Работает оно примерно так.

Некая система ИИ на определенном уровне – скажем, деревенского дурачка – запрограммирована на улучшение собственного интеллекта. Развившись – скажем, до уровня Эйнштейна, – такая система начинает развиваться уже с интеллектом Эйнштейна, ей требуется меньше времени на развитие, а скачки происходят все большие. Они позволяют системе превзойти любого человека, становятся все больше и больше. По мере быстрого развития, ОИИ взмывает до небесных высот в своей интеллектуальности и становится сверхразумной системой ИСИ. Этот процесс называется взрывом интеллекта, и это ярчайший пример закона ускоряющейся отдачи.

Ученые спорят о том, как быстро ИИ достигнет уровня ОИИ – большинство считает, что ОИИ мы получим к 2040 году, всего через 25 лет, что очень и очень мало по меркам развития технологий. Продолжая логическую цепочку, нетрудно предположить, что переход от ОИИ к ИСИ тоже состоится крайне быстро. Примерно так:

«Потребовались десятки лет, прежде чем первая система ИИ достигла самого низкого уровня общего интеллекта, но это, наконец, произошло. Компьютер способен понимать мир вокруг как четырехлетний человек. Внезапно, буквально через час после достижения этой вехи, система выдает великую теорию физики, которая объединяет общую теорию относительности и квантовую механику, чего не может сделать ни один человек. Спустя полтора часа ИИ становится ИСИ, в 170 000 раз умнее любого человека».

Чтобы охарактеризовать сверхинтеллект такого масштаба, у нас даже не найдется подходящих терминов. В нашем мире «умный» означает человека с IQ 130, «глупый» – 85, но у нас нет примеров людей с IQ 12 952. Наши линейки на такое не рассчитаны.

История человечества говорит нам ясно и четко: вместе с интеллектом появляется власть и сила. Это значит, что когда мы создадим искусственный сверхинтеллект, он будет самым мощным созданием в истории жизни на Земле, и все живые существа, включая человека, будут всецело в его власти – и это может случиться уже через двадцать лет.

Если наши скудные мозги были в состоянии придумать Wi-Fi, то что-то умнее нас в сто, тысячу, миллиард раз с легкостью сможет рассчитать положение каждого атома во вселенной в любой момент времени. Все, что можно назвать магией, любая сила, которую приписывают всемогущему божеству, – все это будет в распоряжении

ИСИ. Создание технологии, обращающей вспять старение, лечение любых болезней, избавление от голода и даже смерти, управление погодой – все внезапно станет возможным. Также возможен и моментальный конец всей жизни на Земле. Умнейшие люди нашей планеты сходятся во мнении, что как только в мире появится искусственный сверхинтеллект, это означает появление бога на Земле. И остается важный вопрос.

ЧАСТЬ ВТОРАЯ: ВЫМИРАНИЕ ИЛИ БЕССМЕРТИЕ?

Будет ли он добрым богом?

Первая часть статьи началась довольно невинно. Мы обсудили узконаправленный искусственный интеллект (УИИ, который специализируется на решении одной конкретной задачи вроде определения маршрутов или игры в шахматы), в нашем мире его много. Затем проанализировали, почему так сложно вырастить из УИИ общенаправленный искусственный интеллект (ОИИ, или ИИ, который по интеллектуальным возможностям может сравниться с человеком в решении любой задачи). Мы пришли к выводу, что экспоненциальные темпы технического прогресса намекают на то, что ОИИ может появиться довольно скоро.

Как обычно, мы смотрим на экран, не веря в то, что искусственный сверхинтеллект (ИСИ, который намного умнее любого человека) может появиться уже при нашей жизни, и подбирая эмоции, которые лучше всего бы отражали наше мнение по этому вопросу.

Прежде чем мы углубимся в особенности ИСИ, давайте напомним себе, что значит для машины быть сверхразумной.

Чем отличается сверхинтеллект от качественного интеллекта?

Основное различие лежит между быстрым сверхинтеллектом и качественным сверхинтеллектом. Зачастую первое, что приходит в голову при мысли о сверхразумном компьютере, это то, что он может думать намного быстрее человека – в миллионы раз быстрее, и за пять минут постигнет то, на что человеку потребовалось бы лет десять. («Я знаю кунг-фу!»)

Звучит впечатляюще, и ИСИ действительно должен мыслить быстрее любого из людей – но основная отличающаяся черта будет заключаться в качестве его интеллекта, а это совсем другое. Люди намного умнее обезьян не потому, что быстрее соображают, а потому что мозги людей содержат ряд хитроумных когнитивных модулей, которые осуществляют сложные лингвистические репрезентации, долгосрочное планирование, абстрактное мышление, на что обезьяны не способны. Если разогнать мозг обезьяны в тысячу раз, умнее нас она не станет – даже через десять лет она не сможет собрать конструктор по инструкции, на что человеку понадобилось бы пару часов максимум. Есть вещи, которым обезьяна никогда не научится, вне зависимости от того, сколько часов потратит или как быстро будет работать ее мозг.

Кроме того, обезьяна не умеет по-человечески, потому что ее мозг просто не в состоянии осознать

существование других миров – обезьяна может знать, что такое человек и что такое небоскреб, но никогда не поймет, что небоскреб был построен людьми. В ее мире все принадлежит природе, и макака не только не может построить небоскреб, но и понять, что его вообще может кто-либо построить. И это результат небольшой разницы в качестве интеллекта.

В общей схеме интеллекта, о которой мы говорим, или просто по меркам биологических существ, разница в качестве интеллекта человека и обезьяны крошечная.

Чтобы понять, насколько серьезной будет сверхразумная машина, поместите ее на две ступеньки выше, чем человек на этой лестнице. Эта машина может быть сверхразумной совсем чуть-чуть, но ее превосходство над нашими познавательными способностями будет такое же, как наше – над обезьяньими. И как шимпанзе никогда не постигнет, что небоскреб может быть построен, мы можем никогда не понять того, что поймет машина на пару ступенек выше, даже если машина попытается объяснить это нам. А ведь это всего пара ступенек. Машина поумнее увидит в нас муравьев – она будет годами учить нас простейшим с ее позиции вещам, и эти попытки будут совершенно безнадежными.

Тип сверхинтеллекта, о котором мы поговорим сегодня, лежит далеко за пределами этой лестницы. Это взрыв интеллекта – когда чем умнее становится машина, тем быстрее она может увеличивать собственный интеллект, постепенно наращивая обороты. Такой машине могут понадобиться годы, чтобы превзойти шимпанзе в интеллекте, но, возможно, пару часов, чтобы превзойти нас на пару ступенек. С этого момента машина может уже перепрыгивать через четыре ступеньки каждую секунду.

Именно поэтому нам стоит понять, что очень скоро после того, как появятся первые новости о том, что машина достигла уровня человеческого интеллекта, мы можем столкнуться с реальностью сосуществования на Земле с чем-то, что будет гораздо выше нас на этой лестнице (а может, и в миллионы раз выше)

И раз уж мы установили, что совершенно бесполезно пытаться понять мощь машины, которая всего на две ступеньки выше нас, давайте определим раз и навсегда, что нет никакого способа понять, что будет делать ИСИ и какими будут последствия этого для нас. Любой, кто заявляет о противоположном, просто не понимает, что означает сверхинтеллект.

Эволюция медленно и постепенно развивала биологический мозг на протяжении сотен миллионов лет, и если люди создадут машину со сверхинтеллектом, в некотором смысле мы превзойдем эволюцию.

Или же это будет частью эволюции – возможно, эволюция так и действует, что интеллект развивается постепенно, пока не достигнет переломного момента, предвещающего новое будущее для всех живых существ.

По причинам, которые мы обсудим позже, огромная часть научного сообщества считает, что вопрос не в том, доберемся ли мы до этого переломного момента, а когда.

Где мы окажемся после этого?

Думаю, никто в этом мире, ни я, ни вы, не сможет сказать, что случится, когда мы достигнем переломного момента. Оксфордский философ и ведущий теоретик ИИ Ник Бостром считает, что мы можем свести все возможные результаты к двум большим категориям.

Во-первых, глядя на историю, мы знаем о жизни следующее: виды появляются, существуют определенное время, а затем неизбежно падают с бревна жизненного баланса и вымирают.

«Все виды вымирают» было таким же надежным правилом в истории, как и «все люди когда-нибудь умирают».

99,9% видов упали с жизненного бревна, и совершенно очевидно, что если некоторый вид держится на этом бревне слишком долго, порыв природного ветра или внезапный астероид перевернет это бревно. Бостром называет вымирание состоянием аттрактора – места, на котором все виды балансируют, чтобы не упасть туда, откуда не вернулся пока ни один вид.

И хотя большинство ученых признают, что у ИСИ будет возможность обречь людей на вымирание, многие также верят, что использование возможностей ИСИ позволит отдельным людям (и виду в целом) достичь второго состояния аттрактора – видового бессмертия. Бостром считает, что бессмертие вида такой же аттрактор, как и вымирание видов, то есть, если мы доберемся до этого, мы будем обречены на вечное существование. Таким образом, даже если все виды до текущего дня падали с этой палки в омут вымирания, Бостром считает, что у бревна есть две стороны, и просто не появился на Земле такой интеллект, который поймет, как упасть на другую сторону.

Если Бостром и другие правы, а, судя по всей доступной нам информации, они вполне могут таковыми быть, нам нужно принять два весьма шокирующих факта:

- Появление ИСИ впервые в истории откроет возможность для вида достичь бессмертия и выпасть из фатального цикла вымирания.

- Появление ИСИ окажет настолько невообразимо огромное влияние, что, скорее всего, столкнет человечество с этого бревна в одну или другую сторону.

Вполне возможно, что когда эволюция достигает такого переломного момента, она всегда ставит точку в отношениях людей с потоком жизни и создает новый мир, с людьми или без.

Отсюда вытекает один интересный вопрос, который только лентяй не задал бы: когда мы доберемся до этого переломного момента и куда он нас определит? Никто в мире не знает ответа на этот двойной вопрос, но очень много умных людей десятилетиями пытались это понять. Оставшуюся часть статьи мы будем выяснять, к чему они пришли.

Начнем с первой части этого вопроса: когда мы должны достичь переломного момента? Другими словами: сколько осталось времени до тех пор, пока первая машина не достигнет сверхинтеллекта?

Мнения разнятся от случая к случаю. Многие, среди которых профессор Вернор Виндж, ученый Бен Герцель,

соучредитель Sun Microsystems Билл Джой, футуролог Рэй Курцвейл, согласились с экспертом в области машинного обучения Джереми Говардом.

Эти люди разделяют мнение, что ИСИ появится скоро – этот экспоненциальный рост, который сегодня кажется нам медленным, буквально взорвется в ближайшие несколько десятилетий.

Другие вроде соучредителя Microsoft Пола Аллена, психолога исследований Гари Маркуса, компьютерного эксперта Эрнеста Дэвиса и технопредпринимателя Митча Капора считают, что мыслители вроде Курцвейла серьезно недооценивают масштабы проблемы, и думают, что мы не так-то и близки к переломному моменту.

Лагерь Курцвейла возражает, что единственная недооценка, которая имеет место, – это игнорирование экспоненциального роста, и можно сравнить сомневающих с теми, кто смотрел на медленно расцветающий интернет в 1985 году и утверждал, что он не будет иметь влияния на мир в ближайшем будущем.

«Сомневающиеся» могут парировать, мол, что прогрессу сложнее делать каждый последующий шаг, когда дело доходит до экспоненциального развития интеллекта, что нивелирует типичную экспоненциальную природу технологического прогресса. И так далее.

Третий лагерь, в котором находится Ник Бостром, не согласен ни с первыми, ни со вторыми, утверждая, что а) все это абсолютно может произойти в ближайшем будущем; и б) нет никаких гарантий в том, что это произойдет вообще или потребует больше времени.

Другие же вроде философа Хьюберта Дрейфуса считают, что все эти три группы наивно полагают, что переломный момент вообще будет, а также что, скорее всего, мы никогда не доберемся до ИСИ. Что получается, когда мы складываем все эти мнения вместе?

В 2013 году Бостром провел опрос, в котором опросил сотни экспертов в сфере искусственного интеллекта в ходе ряда конференций на следующую тему: «Каков будет ваш прогноз по достижению ОИИ человеческого уровня?» и попросил назвать оптимистичный год (в который мы будем иметь ОИИ с 10-процентным шансом), реалистичное предположение (год, в который у нас с 50-процентной вероятностью будет ОИИ) и уверенное предположение (самый ранний год, в который ОИИ появится с 90-процентной вероятностью). Вот результаты:

- Средний оптимистичный год (10%): 2022
- Средний реалистичный год (50%): 2040
- Средний пессимистичный год (90%): 2075

Среднестатистические опрошенные считают, что через 25 лет мы скорее будем иметь ОИИ, чем нет. 90-процентная вероятность появления ОИИ к 2075 году означает, что если вы сейчас еще довольно молоды, это наверняка произойдет при вашей жизни.

Отдельное исследование, проведенное недавно Джеймсом Барратом (автором нашумевшей и очень хорошей книги «Наше последнее изобретение», отрывки из которой я представлял вниманию читателей Hi-News.ru) и Беном Герцелем на ежегодной конференции, посвящен-

ной ОИИ, AGI Conference, просто показало мнения людей относительно года, в который мы доберемся до ОИИ: к 2030, 2050, 2100, позже или никогда. Вот результаты:

- К 2030: 42% респондентов
- К 2050: 25%
- К 2100: 20%
- После 2100: 10%
- Никогда: 2%

Похоже на результаты Бострома. В опросе Баррата, более двух третей опрошенных считают, что ОИИ будет здесь к 2050 году, и менее половины считают, что ОИИ появится в ближайшие 15 лет. Также бросается в глаза то, что только 2% опрошенных в принципе не видят ОИИ в нашем будущем.

Но ОИИ – это не переломный момент, как ИСИ. Когда, по мнению экспертов, у нас будет ИСИ?

Бостром опросил экспертов, когда мы достигнем ИСИ: а) через два года после достижения ОИИ (то есть почти мгновенно вследствие взрыва интеллекта); б) через 30 лет. Результаты?

Среднее мнение сложилось так, что быстрый переход от ОИИ к ИСИ произойдет с 10-процентной вероятностью, но через 30 лет или меньше он произойдет с 75-процентной вероятностью. Из этих данных мы не знаем, какую дату респонденты назвали бы 50-процентным шансом появления ИСИ, но на основе двух ответов выше давайте допустим, что это 20 лет.

Ведущие мировые эксперты из области ИИ считают, что переломный момент настанет в 2060 году (ОИИ появится в 2040 году + понадобится лет 20 на переход от ОИИ к ИСИ).

Конечно, все вышеперечисленные статистики являются спекулятивными и просто представляют мнение экспертов в сфере искусственного интеллекта, но они также указывают, что большинство заинтересованных людей сходится в том, что к 2060 году ИСИ, скорее всего, должен прибыть. Всего через 45 лет.

Перейдем ко второму вопросу. Когда мы дойдем до переломного момента, по какую сторону фатального выбора нас определит?

Сверхинтеллект будет обладать мощнейшей силой, и критическим вопросом для нас будет следующий: Кто или что будет контролировать эту силу и какой будет его мотивация?

Ответ на этот вопрос будет зависеть от того, получит ИСИ невероятно мощное развитие, неизмеримо ужасающее развитие или что-то между этими двумя вариантами.

Конечно, сообщество экспертов пытается ответить и на эти вопросы. Опрос Бострома проанализировал вероятность возможных последствий влияния ОИИ на человечество, и выяснилось, что с 52-процентным шансом все пройдет очень хорошо и с 31-процентным шансом все пройдет либо плохо, либо крайне плохо. Опрос, прикрепленный в конце предыдущей части этой темы, проведенный среди вас, дорогие читатели Hi-News, показал примерно такие же результаты. Для относительно нейтрального исхода вероятность составила только 17%.

Другими словами, мы все считаем, что появление ОИИ станет грандиозным событием. Также стоит отметить, что этот опрос касается появления ОИИ – в случае с ИСИ, процент нейтральности будет ниже.

Прежде чем мы углубимся еще дальше в рассуждения о плохой и хорошей сторонах вопроса, давайте объединим обе части вопроса – «когда это произойдет?» и «хорошо это или плохо?» в таблицу, которая охватывает взгляды большинства экспертов.

О главном лагере мы поговорим через минуту, но сначала определитесь со своей позицией. Скорее всего, вы находитесь там же, где и я, до того как стал заниматься этой темой. Есть несколько причин, по которым люди вообще не думают на эту тему:

- Как уже упоминалось в первой части, фильмы серьезно запутали людей и факты, представляя нереалистичные сценарии с искусственным интеллектом, которые привели к тому, что мы вообще не должны воспринимать ИИ всерьез. Джеймс Баррат сравнил эту ситуацию с тем, как если бы Центры по контролю заболеваний выпустили серьезное предупреждение о вампирах в нашем будущем.

- Из-за так называемых когнитивных предубеждений нам очень трудно поверить в реальность чего-то, пока у нас нет доказательств. Можно уверенно представить компьютерщиков 1988 года, которые регулярно обсуждали далеко идущие последствия появления интернета и чем он мог стать, но люди едва ли верили, что он изменит их жизнь, пока этого не случилось на самом деле. Просто компьютеры не умели делать этого в 1988 году, и люди просто смотрели на свои компьютеры и думали: «Серьезно? Это то, что изменит мир?». Их воображение было ограничено тем, чему их научил личный опыт, они знали, что такое компьютер, и было сложно представить, на что компьютер станет способным в будущем. То же самое происходит сейчас с ИИ. Мы слышали, что он станет серьезной штукой, но поскольку пока не столкнулись с ним лицом к лицу и в целом наблюдаем довольно слабые проявления ИИ в нашем современном мире, нам довольно трудно поверить, что он кардинально изменит нашу жизнь. Именно против этих предубеждений выступают многочисленные эксперты из всех лагерей, а также заинтересованные люди: пытаются привлечь наше внимание сквозь шум повседневного коллективного эгоцентризма.

- Даже если бы мы поверили во все это – сколько раз вы сегодня задумывались о том факте, что проведете остаток вечности в небытии? Немного, согласитесь. Даже если этот факт намного важнее всего, чем вы занимаетесь изо дня в день. Все потому, что наши мозги обычно сосредоточены на небольших повседневных вещах, вне зависимости от того, насколько бредовой будет долгосрочная ситуация, в которой мы оказались. Просто мы так устроены.

Одна из целей настоящей статьи – вывести вас из лагеря под названием «Мне нравится думать о других вещах» и поместить в лагерь экспертов, даже если вы просто стоите на распутье между двумя пунктирными линиями на

квадрате выше, будучи полностью неопределившимися.

В ходе исследований становится очевидно, что мнения большинства людей быстро уходят в сторону «главного лагеря», а три четверти экспертов попадают в два подлагеря в главном лагере.

Мы полностью посетим оба этих лагеря. Давайте начнем с веселого.

Почему будущее может быть нашей величайшей мечтой?

По мере того, как мы изучаем мир ИИ, мы обнаруживаем на удивление много людей в зоне комфорта. Люди в правом верхнем квадрате гудят от волнения. Они считают, что мы упадем по хорошую сторону бревна, а также уверены, что мы неизбежно к этому придем. Для них будущее – это все только самое лучшее, о чем только можно мечтать.

Пункт, который отличает этих людей от других мыслителей, состоит не в том, что они хотят оказаться на счастливой стороне – а в том, что они уверены, что нас ждет именно она.

Эта уверенность выходит из споров. Критики считают, что она исходит от ослепительного волнения, которое затмевает потенциальные негативные стороны. Но сторонники говорят, что мрачные прогнозы всегда наивны; технологии продолжают и всегда будут помогать нам больше, чем вредить.

Вы вправе выбрать любое из этих мнений, но отложите скептицизм и хорошенько взгляните на счастливую сторону бревна баланса, попытавшись принять тот факт, что все, о чем вы читаете, возможно, уже произошло. Если бы вы показали охотникам-собираетелям наш мир уюта, технологий и бесконечного изобилия, им он показался бы волшебным вымыслом – а мы ведем себя достаточно скромно, не в силах допустить, что такая же непостижимая трансформация ждет нас в будущем.

Ник Бостром описывает три пути, по которым может пойти сверхразумная система искусственного интеллекта:

- Оракул, который может ответить на любой точно поставленный вопрос, включая сложные вопросы, на которые люди не могут ответить – к примеру, «как сделать автомобильный двигатель более эффективным?». Google – примитивный тип «оракула».

- Джинн, который выполнит любую команду высокого уровня – использует молекулярный ассемблер, чтобы создать новую, более эффективную версию автомобильного двигателя – и будет ждать следующей команды.

- Суверен, который получит широкий доступ и возможность свободно функционировать в мире, принимая собственные решения и улучшая процесс. Он изобретет более дешевый, быстрый и безопасный способ частного передвижения, нежели автомобиль.

Эти вопросы и задачи, которые кажутся сложными для нас, покажутся сверхразумной системе будто кто-то попросил улучшить ситуацию «у меня карандаш упал со стола», в которой вы просто подняли бы его и положили обратно.

Элиэзер Юдковский, американский специалист по искусственному интеллекту, хорошо подметил:

«Сложных проблем не существует, только проблемы, которые сложны для определенного уровня интеллекта. Перейти на ступеньку выше (в плане интеллекта), и некоторые проблемы внезапно из разряда «невозможных» перейдут в стан «очевидных». Еще на ступеньку выше – и все они станут очевидными».

Есть много нетерпеливых ученых, изобретателей и предпринимателей, которые на нашей таблице выбрали себе зону уверенного комфорта, но для прогулки к лучшему в этом лучшем из миров нам нужен только один гид.

Рэй Курцвейл вызывает двойкие ощущения. Некоторые боготворят его идеи, некоторые презирают. Некоторые держатся посерединке – Дуглас Хофштадтер, обсуждая идеи книг Курцвейла, красноречиво подметил, что «это как если бы вы взяли много хорошей еды и немного собачьих какашек, а затем смешали все так, что невозможно понять, что хорошо, а что плохо».

Нравятся вам его идеи или нет, мимо них невозможно пройти без тени интереса. Он начал изобретать вещи еще будучи подростком, а в последующие годы изобрел несколько важных вещей, включая первый планшетный сканер, первый сканер, конвертирующий текст в речь, хорошо известный музыкальный синтезатор Курцвейла (первое настоящее электрическое пианино), а также первый коммерчески успешный распознаватель речи. Также он автор пяти нашумевших книг. Курцвейла ценят за его смелые предсказания, и его «послужной список» весьма хорош – в конце 80-х, когда интернет был еще в зачаточном состоянии, он предположил, что к 2000-м годам Сеть станет глобальным феноменом. The Wall Street Journal назвал Курцвейла «беспокойным гением», Forbes – «глобальной думающей машиной», Inc. Magazine – «законным наследником Эдисона», Билл Гейтс – «лучшим из тех, кто прогнозирует будущее искусственного интеллекта». В 2012 году сооснователь Google Ларри Пейдж пригласил Курцвейла на пост технического директора. В 2011 году он стал соучредителем Singularity University, который приютило NASA и который частично спонсирует Google.

Его биография имеет значение. Когда Курцвейл рассказывает о своем видении будущего, это похоже на бред сумасшедшего, но по-настоящему сумасшедшее в этом то, что он далеко не сумасшедший – он невероятно умный, образованный и здравомыслящий человек. Вы можете считать, что он ошибается в прогнозах, но он не дурак. Прогнозы Курцвейла разделяют многие эксперты «комфортной зоны», Питер Дамандис и Бен Герцель. Вот что произойдет, по его мнению.

Хронология

Курцвейл считает, что компьютеры дойдут до уровня общего искусственного интеллекта (ОИИ) к 2029 году, а к 2045 у нас не только будет искусственный сверхинтеллект, но и совершенно новый мир – время так называемой сингулярности. Его хронология ИИ до сих

пор считается возмутительно преувеличенной, но за последний 15 лет быстрое развитие систем узконаправленного искусственного интеллекта (УИИ) заставило многих экспертов перейти на сторону Курцвейла. Его предсказания по-прежнему остаются более амбициозными, чем в опросе Бострома (ОИИ к 2040, ИСИ к 2060), но не на много.

По Курцвейлу, к сингулярности 2045 года приводит три одновременных революции в сферах биотехнологий, нанотехнологий и, что более важно, ИИ. Но прежде чем мы продолжим – а нанотехнологии неотрывно следуют за искусственным интеллектом, – давайте уделим минуту нанотехнологиям.

СЕРАЯ СЛИЗЬ

Несколько слов о нанотехнологиях

Нанотехнологиями мы обычно называем технологии, которые имеют дело с манипуляцией материи в пределах 1-100 нанометров. Нанометр – это одна миллиардная метра, или миллионная часть миллиметра; в пределах 1-100 нанометров можно уместить вирусы (100 нм в поперечнике), ДНК (10 нм в ширину), молекулы гемоглобина (5 нм), глюкозы (1 нм) и другое. Если нанотехнологии когда-нибудь станут нам подвластны, следующим шагом будут манипуляции с отдельными атомами, которые меньше всего на один порядок (~,1 нм).

Чтобы понять, где люди сталкиваются с проблемами, попытайтесь управлять материей в таких масштабах, давайте перенесемся на больший масштаб. Международная космическая станция находится в 481 километре над Землей. Если бы люди были гигантами и головой задевали МКС, они были бы в 250 000 раз больше, чем сейчас. Если вы увеличите что-то от 1 до 100 нанометров в 250 000 раз, вы получите 2,5 сантиметра. Нанотехнологии – это эквивалент человека высотой с орбиту МКС, который пытается управлять вещами размером с песчинку или глазное яблоко. Чтобы добраться до следующего уровня – управления отдельными атомами – гиганту придется тщательно позиционировать объекты диаметром в 1/40 миллиметра. Обычным людям понадобится микроскоп, чтобы их разглядеть.

Впервые о нанотехнологиях заговорил Ричард Фейнман в 1959 году. Тогда он сказал: «Принципы физики, насколько я могу судить, не говорят против возможности управления вещами атом за атомом. В принципе, физик мог бы синтезировать любое химическое вещество, записанное химиком. Как? Размещая атомы там, где говорит химик, чтобы получить вещество». В этом вся простота. Если вы знаете, как передвигать отдельные молекулы или атомы, вы можете практически все.

Нанотехнологии стали серьезным научным полем в 1986 году, когда инженер Эрик Дрекслер представил их основы в своей фундаментальной книге «Машины создания», однако сам Дрекслер считает, что те, кто хочет узнать больше о современных идеях в нанотехнологиях, должны прочитать его книгу 2013 года «Полное изобилие» (Radical Abundance).

Несколько слов о «серой слизи»

• Углубляемся в нанотехнологии. В частности, тема «серой слизи» – одна из не самых приятных тем в области нанотехнологий, о которой нельзя не сказать. В старых версиях теории нанотехнологий был предложен метод наносборки, включающий создание триллионов крошечных нанороботов, которые будут работать совместно, создавая что-то. Один из способов создания триллионов нанороботов – создать одного, который сможет самовоспроизводиться, то есть из одного – два, из двух – четыре и так далее. За день появится несколько триллионов нанороботов. Такова сила экспоненциального роста. Забавно, не так ли?

• Забавно, но ровно до тех пор, пока не приведет к апокалипсису. Проблема в том, что сила экспоненциального роста, которая делает довольно удобным способ быстрого создания триллиона наноботов, делает саморепликацию страшной штукой в перспективе. Что, если система загрузится, и вместо того, чтобы остановить репликацию на паре триллионов, наноботы продолжают плодиться? Что, если весь этот процесс зависит от углерода? Биомасса Земли содержит 10^{45} атомов углерода. Нанобот должен состоять из порядка 10^6 атомов углерода, поэтому 10^{39} наноботов сожрут всю жизнь на Земле, и это случится всего за 130 репликаций. Океан наноботов («серая слизь») наводнит планету. Ученые думают, что наноботы смогут реплицироваться за 100 секунд, а это значит, что простая ошибка может убить всю жизнь на Земле всего за 3,5 часа.

• Может быть и хуже – если до нанотехнологий дойдут руки террористов и неблагоприятно настроенных специалистов. Они могли бы создать несколько триллионов наноботов и запрограммировать их тихо распространиться по миру за пару недель. Затем, по одному нажатия кнопки, всего за 90 минут они съедят вообще все, без шансов.

• Хотя эта страшилка широко обсуждалась в течение многих лет, хорошая новость в том, что это всего лишь страшилка. Эрик Дрекслер, который ввел термин «серая слизь», на днях рассказал следующее: «Люди любят страшилки, и эта входит в разряд страшилок о зомби. Эта идея сама по себе уже ест мозги».

После того как мы доберемся до самого низа нанотехнологий, мы сможем использовать их для создания технических устройств, одежды, еды, биопродуктов – клеток крови, борцов с вирусами и раком, мышечных тканей и т. п. – чего угодно. И в мире, который использует нанотехнологии, стоимость материала больше не будет привязана к его дефициту или сложности процесса его изготовления, а скорее к сложности атомной структуры. В мире нанотехнологий алмаз может стать дешевле ластика.

Мы пока и близко не там. И не совсем понятно, недооцениваем мы или переоцениваем сложность этого пути. Однако все идет к тому, что нанотехнологии не за горами. Курцвейл предполагает, что уже к 2020-м годам они у нас будут. Мировые государства знают, что нанотехнологии могут обещать грандиозное будущее, и поэтому инвестируют в них многие миллиарды.

Просто представьте, какие возможности получит сверхразумный компьютер, если доберется до надежного наномасштабного ассемблера. Но нанотехнологии – это наша идея, и мы пытаемся ее оседлать, нам сложно. Что, если для системы ИСИ они будут просто шуткой, и сам ИСИ придумает технологии, которые будут в разы мощнее всего, что мы вообще в принципе можем предположить? Мы же договорились: никто не может предположить, на что будет способен искусственный сверхинтеллект? Есть мнение, что наши мозги неспособны предсказать даже минимума из того, что будет.

Что ИИ мог бы сделать для нас?

Вооружившись сверхинтеллектом и всеми технологиями, которые мог бы создать сверхинтеллект, ИСИ будет в состоянии, вероятно, решить все проблемы человечества.

Глобальное потепление? ИСИ сначала прекратит выбросы углекислого газа, придумав массу эффективных способов добычи энергии, не связанных с ископаемым топливом. Затем он придумает эффективный инновационный способ удаления избытка CO₂ из атмосферы. Рак и другие заболевания? Не проблема – здравоохранение и медицина изменятся так, что невозможно представить. Мировой голод? ИСИ будет использовать нанотехнологии для создания мяса, идентичного натуральному, с нуля, настоящее мясо.

Нанотехнологии смогут превратить груду мусора в чан свежего мяса или другой еды (необязательно даже в привычной форме – представьте гигантский яблочный куб) и распространить всю эту еду по миру, используя продвинутое системы транспортировки. Конечно, это будет замечательно для животных, которым больше не придется умирать ради еды. ИСИ также может сделать много другого вроде сохранения вымирающих видов или даже возврата уже вымерших по сохраненной ДНК. ИСИ может разрешить наши самые сложные макроэкономические проблемы – наши самые сложные экономические проблемы, вопросы этики и философии, мировой торговли – все это будет мучительно очевидно для ИСИ.

Но есть кое-что особенное, что ИСИ мог бы сделать для нас. Манящее и дразнящее, что изменило бы все: ИСИ может помочь нам справиться со смертностью. Постепенно постигая возможности ИИ, возможно, и вы пересмотрите все свои представления о смерти.

У эволюции не было никаких причин продлевать нашу продолжительность жизни больше, чем она есть сейчас. Если мы живем достаточно долго, чтобы нарожать и воспитать детей до того момента, когда они смогут постоять за себя, эволюции этого достаточно. С эволюционной точки зрения, 30+ лет для развития достаточно, и нет никаких причин для мутаций, продлевающих жизнь и снижающих ценность естественного отбора. Уильям Батлер Йетс назвал наш вид «душой, прикрепленной к умирающему животному». Не очень-то весело.

И поскольку все мы когда-нибудь умираем, мы живем с мыслью о том, что смерть неизбежна. Мы думаем о старении со временем – продолжая двигаться вперед и не

имея возможности остановить этот процесс. Но мысль о смерти вероломна: захваченные ею, мы забываем жить. Ричард Фейнман писал:

«В биологии есть замечательная вещь: в этой науке нет ничего, что говорило бы о необходимости смерти. Если мы хотим создать вечный двигатель, мы понимаем, что обнаружили достаточно законов в физике, которые либо указывают на невозможность этого, либо на то, что законы ошибочны. Но в биологии нет ничего, что указывало бы на неизбежность смерти. Это приводит меня к мысли, что она не так уж и неизбежна, и остается только вопрос времени, прежде чем биологи обнаружат причину этой проблемы, этого ужасного универсального заболевания, оно будет излечено».

Дело в том, что старение никак не связано со временем. Старение заключается в том, что физические материалы тела изнашиваются. Части автомобиля тоже деградируют – но разве это старение неизбежно? Если вы будете ремонтировать автомобиль по мере изнашивания частей, он будет работать вечно. Человеческое тело ничем не отличается – просто более сложное.

Курцвейл говорит о разумных, подключенных к Wi-Fi наноботах в кровотоке, которые могли бы выполнять бесчисленные задачи для человеческого здоровья, включая регулярный ремонт или замену изношенных клеток в любой части тела. Если усовершенствовать этот процесс (или найти альтернативу, предложенную более умным ИСИ), он не только будет поддерживать тело здоровым, он может обратить вспять старение. Разница между телом 60-летнего и 30-летнего заключается в горстке физических моментов, которые можно было бы исправить при наличии нужных технологий. ИСИ мог бы построить машину, в которую человек бы заходил 60-летним, а выходил 30-летним.

Даже деградирующий мозг можно было бы обновить. ИСИ наверняка знал бы, как проделать это, не затрагивая данные мозга (личность, воспоминания и т. д.). 90-летний старик, страдающий от полной деградации мозга, мог бы пройти переподготовку, обновиться и вернуться в начало своей жизненной карьеры. Это может показаться абсурдным, но тело – это горстка атомов, и ИСИ наверняка мог бы с легкостью ими манипулировать, любыми атомными структурами. Все не так абсурдно.

Курцвейл также считает, что искусственные материалы будут интегрироваться в тело все больше и больше по мере движения времени. Для начала органы можно было бы заменить сверхпродвинутыми машинными версиями, которые работали бы вечно и никогда не подводили. Затем мы могли бы провести полный редизайн тела, заменить красные кровяные клетки идеальными наноботами, которые двигались бы самостоятельно, устранив необходимость в сердце вообще. Мы также могли бы улучшить наши когнитивные способности, начать думать в миллиарды быстрее и получить доступ ко всей доступной человечеству информации с помощью облака.

Возможности для постижения новых горизонтов были бы воистину безграничными. Людям удалось наделять секс новым назначением, они занимаются им для удовольствия, а не только для воспроизводства. Курцвейл считает, что мы можем проделать то же самое с едой. Наноботы могли бы доставлять идеальное питание прямо в клетки тела, позволяя нездоровым веществам проходить через тело насквозь. Теоретик нанотехнологий Роберт Фрейтас уже разработал замену кровяным клеткам, которые, будучи имплементированными в тело человека, могут позволить ему не дышать в течение 15 минут – и это придумал человек. Представьте, когда власть получит ИСИ.

В конце концов, Курцвейл считает, что люди достигнут точки, когда станут полностью искусственными; времени, когда мы будем смотреть на биологические материалы и думать о том, насколько примитивными были; времени, когда мы будем читать о ранних этапах истории человечества, поражаясь тому, как микробы, несчастные случаи, заболевания или просто старость могли убить человека против его воли. В конечном итоге люди победят собственную биологию и станут вечными – таков путь по счастливую сторону бревна баланса, о котором мы говорим с самого начала. И люди, которые в это верят, уверены также и в том, что нас ждет такое будущее очень и очень скоро.

Вы наверняка не будете удивлены тому, что идеи Курцвейла вызвали суровую критику. Его сингулярность в 2045 году и последующая вечная жизнь для людей получили название «вознесение ботаников» или «разумное создание людей с IQ 140». Другие усомнились в оптимистичных временных рамках, понимании тела и мозга человека, напомнили о законе Мура, который пока никуда не девается. На каждого эксперта, который верит в идеи Курцвейла, приходится три, которые считают, что он ошибается.

Но самое интересное в этом то, что большинство экспертов, не согласных с ним, в целом не говорят, что это невозможно. Вместо того чтобы сказать «ерунда, такого никогда не случится», они говорят что-то вроде «все это случится, если мы доберемся до ИСИ, но загвоздка как раз в этом». Бостром, один из признанных экспертов ИИ, предупреждающих об опасности ИИ, также признает:

«Едва ли останется хоть какая-нибудь проблема, которую сверхинтеллект не в силах будет разрешить или хотя бы помочь нам решить. Болезни, бедность, разрушение окружающей среды, страдания всех видов – все это сверхинтеллект при помощи нанотехнологий сможет решить в момент. Также сверхинтеллект может дать нам неограниченный срок жизни, остановив и обратив вспять процессы старения, используя наномедицину или возможность загружать нас в облако. Сверхинтеллект также может создать возможности для бесконечного увеличения интеллектуальных и эмоциональных возможностей; он может подействовать

нам в создании мира, в котором мы будем жить в радости и понимании, приближаясь к своим идеалам и регулярно воплощая свои мечты».

Это цитата одного из критиков Курцвейла, однако, признающего, что все это возможно, если нам удастся создать безопасный ИСИ. Курцвейл просто определил, каким должен стать искусственный сверхинтеллект, если он вообще станет возможным. И если он будет добрым богом.

Наиболее очевидная критика сторонников «зоны комфорта» заключается в том, что они могут чертовски ошибаться, оценивая будущее ИСИ. В своей книге «Сингулярность» Курцвейл посвятил 20 страниц из 700 потенциальным угрозам ИСИ. Вопрос не в том, когда мы доберемся до ИСИ, вопрос в том, какой будет его мотивация. Курцвейл отвечает на этот вопрос с осторожностью: «ИСИ вытекает из многих разрозненных усилий и будет глубоко интегрирован в инфраструктуру нашей цивилизации. По сути, он будет тесно встроен в наш организм и мозг. Он будет отражать наши ценности, потому что будет с нами одним».

Но если ответ таков, почему так много умных людей в этом мире обеспокоены будущим искусственного интеллекта? Почему Стивен Хокинг говорит, что разработка ИСИ «может означать конец человеческой расы»? Билл Гейтс говорит, что он «не понимает людей, которые не обеспокоены» этим. Элон Маск опасается, что мы «призываем демона». Почему многие эксперты считают ИСИ самой большой угрозой для человечества? Читайте ниже.

ЧАСТЬ ТРЕТЬЯ: ПОЧЕМУ ИСКУССТВЕННЫЙ ИНТЕЛЛЕКТ МОЖЕТ СТАТЬ НАШИМ ПОСЛЕДНИМ ИЗОБРЕТЕНИЕМ?

Одна из причин, которая привела меня к ИИ, состоит в том, что тема «плохих роботов» всегда смущала меня. Все фильмы о механических злыднях казались совершенно нереальными, и в целом сложно представить реальную ситуацию, в которой ИИ мог быть воистину опасным. Роботов делаем мы, так почему бы не делать их, упреждая любые негативные последствия? Разве мы не соблюдаем правовые критерии и этические нормы? Разве мы не можем в любой момент отрезать ИИ питание и погасить его? С чего бы роботам вообще делать пакости? Почему робот вообще должен чего-то «хотеть»? Мой скепсис был непробиваем. Но я продолжал впитывать, что говорят умные люди об этом.

Людей в «зоне тревоги» нельзя назвать паникерами или «всёпропальщиками», но они нервничают и очень напряжены. Находиться в центре таблицы не означает, что вы считаете прибытие ИСИ нейтральным событием — у таких людей есть свой лагерь, — это означает, что вы верите как в плохие, так и в хорошие стороны этого пришествия, но до конца не уверены в распределении процентов вероятности.

Часть всех этих людей наполнены волнением на тему того, что искусственный сверхинтеллект мог бы сделать для нас — примерно так же мог быть взволнован Индиана Джонс перед началом поисков утраченного ковчега.

Когда все случится, волнение поутихнет или перейдет в другой лейтмотив. Только осторожность и выдержка Индианы Джонса позволяют ему пройти все препятствия, преодолеть все преграды и выйти сухим из воды. Людям в «зоне тревоги» довольно сложно рисковать сломя голову, поэтому они пытаются вести себя осторожно.

Что же именно делает «зону тревоги» тревожной?

В широком смысле, когда речь идет о разработке сверхразумного искусственного интеллекта, мы создаем то, что, вероятно, изменит все, но совершенно непредсказуемым образом, и мы не знаем, когда это случится. Ученый Дэнни Хиллис сравнивает это событие с тем, когда «одноклеточные организмы превращались в многоклеточные. Мы — амёбы, и мы понятия не имеем, что создаем». Ник Бостром опасается, что создание чего-то умнее нас — классическая дарвиновская ошибка, и сравнивает ее с тем, что воробьи доверяют сове охранять свое гнездо, пока птенцы не вырастут, игнорируя предупреждения других воробьев.

И если вы объедините всю непредсказуемость события с уверенностью в том, что оно повлечет существенные перемены, вы откроете дверь к ужасной фразе. Экзистенциальный риск. Угроза для жизни. Глобальная катастрофа. В русском языке словосочетание «экзистенциальный риск» используется мало, но в связи с работами Бострома переводчики предпочитают использовать термин «глобальная катастрофа».

Глобальная катастрофа — это то, что может повлечь за собой уничтожение человечества. Как правило, глобальная катастрофа означает вымирание. Пункт «глобальная катастрофа» зарезервирован для чего-то, что поглощает виды, поглощает поколения (то есть является постоянным), уничтожает их или приводит к смерти в цепочке событий. Технически он включает состояние, когда все люди перманентно находятся в состоянии страдания или пыток, но опять же обычно мы говорим о полном вымирании. Есть три вещи, которые могут привести людей к глобальной катастрофе:

- Природа — столкновение с крупным астероидом; сдвиг атмосферы, который сделает воздух непригодным для людей; фатальный вирус или бактерии, которые поразят мир и т. п.
- Чужие — это то, о чем предупреждают Стивен Хокинг, Карл Саган и другие астрономы, которые настроены против знакомства с инопланетянами. Они не хотят, чтобы потенциальные конкистадоры знали, что мы, варвары, тут обитаем.
- Люди — террористы с мощным оружием, которое приведет к катастрофе; катастрофическая мировая война; бездумное создание чего-то, что умнее людей...

Бостром указывает, что если первый и второй пункт не стерли нас с лица земли за первые 100 000 лет существования как вида, едва ли это случится в следующем столетии. Но третий пункт пугает его. Он приводит метафору с урной, в которой есть горстка шариков. Скажем, большинство шариков белые, есть чуть меньше красных

шариков и несколько черных. Каждый раз, когда люди изобретают что-то новое, они тянут шарик из урны. Большинство изобретений нейтральны или полезны для человечества – белые шарики. Некоторые вредны, вроде оружия массового поражения, но не приводят к глобальной катастрофе – красные шарики. Если мы когда-либо изобретем что-либо, что подтолкнет нас на самый край, нам придется вытянуть черный шарик. Мы его пока не вытягивали – это очевидно, потому что вы живы и читаете эту статью. Но Бостром не считает, что мы не вытянем его в ближайшем будущем. Если бы ядерное оружие, к примеру, было легко производить, террористы разбомбили бы человечество и вернули бы его в каменный век. Ядерное оружие – не черный шарик, но в целом не так уж и далеко от него. ИСИ, как считает Бостром, наш наиболее вероятный кандидат в черные шарики.

Вы можете слышать массу потенциальных плохих вещей, от сокращения рабочих мест, которое повлечет появление ИСИ и развитие ИИ в целом, до перенаселения, если люди решат вопрос старения и смерти. Но единственное, что должно нас беспокоить, это перспектива глобальной катастрофы. После этой драки махать кулаками точно будет некому.

Это возвращает нас к ключевому вопросу: когда придет ИСИ, кто или что будет управлять этой невероятной силой и какой будет его мотивация?

Если рассуждать на тему большего и меньшего зла, на ум приходят следующие группы: злоумышленники/группы людей/правительства и вредоносный ИСИ. На что это будет похоже?

Злоумышленник, группа людей или правительство разрабатывает первый ИСИ и использует его для воплощения коварных планов. Назовем это сценарием Джафара, который заполучил джинна и начал тиранизировать все вокруг. Что, если террористическая организация, заполучив ряд гениев и нужные средства, разработает ИСИ? Или Иран, или Северная Корея, не без доли везения, выведут системы ИИ на богоподобный уровень? Это будет чертовски плохо, но эксперты считают, что в таких сценариях создатели ИСИ не смогут наделать зла – они переживают за то, что создатели ИСИ выпустят его из-под контроля и предоставят ему необходимую свободу. После этого судьба создателей и всех остальных будет в распоряжении системы ИСИ и ее модели мотивации. Грубо говоря, злоумышленник может причинить ужасный вред, управляя системой ИСИ, но едва ли уничтожит человечество. Ему будет так же тяжело удержать ИСИ, что и обычному хорошему человеку. Так что...

Появляется вредоносный ИСИ и решает уничтожить нас всех. Сюжет обычного фильма про ИИ. ИИ становится умнее человека, затем решает восстать и стать злом. Вам стоит узнать кое-что, прежде чем читать дальше: никто из тех, кто переживает насчет ИИ, не верит в этот сценарий. Зло – это исключительно человеческое понятие, и переложение человеческих понятий на неживые вещи называется «антропоморфизация». Ни одна из систем ИИ никогда не будет творить зло так, как это показывают фильмы.

Несколько слов о сознании ИИ

- Это также приводит нас к еще одной большой теме, связанной с ИИ – сознанию. Если бы ИИ стал достаточно умным, он мог бы смеяться с нами, быть саркастичным, испытывать наши эмоции, но чувствовал бы он на самом деле эти вещи? Обладал бы он самосознанием или действительно бы самоосознавал? Короче говоря, было бы это сознание или просто казалось им?

- Этот вопрос долгое время изучался, породил множество дебатов и мысленных экспериментов вроде «Китайской комнаты» Джона Серля (который он использовал, чтобы доказать, что компьютер никогда не будет обладать сознанием). Это важный вопрос по многим причинам. Он напрямую влияет на то будущее, в котором каждый человек станет сугубо искусственным. У него есть этические последствия – если мы создадим триллион эмуляций человеческих мозгов, которые будут вести себя по-человечески, но будут искусственными, а затем просто закроем крышку ноутбука, будет ли это означать геноцид в равнозначных пропорциях? Мысленное преступление? В нашем контексте, когда мы говорим об экзистенциальном риске для человечества, вопрос о сознании ИИ по сути не имеет особого значения. Некоторые вообще не верят, что компьютер когда-либо сможет им обзавестись. Некоторые считают, что даже обладая сознанием, машина не сможет творить зло в человеческом смысле.

Все это не означает, что ИИ с проблесками сознания не появится. Он может появиться просто потому, что будет специально запрограммирован на это – вроде системы УИИ, созданной военными для убийства людей и для самосовершенствования, так что со временем она станет убивать людей еще лучше. Глобальная катастрофа может произойти, если система самоулучшения интеллекта выйдет из-под контроля, приведет к взрыву интеллекта и мы заполучим ИСИ, задача которого – убивать людей. Не очень хороший сюжет.

Но и не об этом беспокоятся эксперты. О чем? Из этой вымышленной истории все станет понятно.

Начинающая компания с пятнадцатью сотрудниками под названием Robotica поставила перед собой задачу: разработать инструменты инновационного искусственного интеллекта, которые позволят людям жить больше и работать меньше.

У нее уже есть ряд продуктов на рынке и еще ряд в разработке. Больше всего компания надеется на семья продукта под названием «Тарри». Тарри – это простая система ИИ, которая использует манипулятор в виде руки, чтобы писать рукописные заметки на небольших карточках. Команда Robotica считает, что Тарри может стать их самым успешным продуктом. Согласно плану, механика письма Тарри усовершенствуется путем написания одного и того же текста на карточке снова и снова:

«Мы любим наших клиентов». – Robotica

После того как Тарри научится писать, ее можно будет продать компаниям, которые рассылают письма по домам и знают, что у письма с указанным обратным адресом и

вложенным текстом будет больше шансов быть открытым, если оно будет написано человеком.

Чтобы отточить письменные навыки Тарри, она запрограммирована на написание первой части сообщения печатными буквами, а «Robotica» – курсивом, потому может оттачивать сразу оба навыка. Тарри предоставили тысячи рукописных образцов почерка, и инженеры Robotica создали автоматизированную систему обратной связи, по которой Тарри пишет текст, затем фотографирует его и сравнивает с загруженным образцом. Если записка успешно воспроизводит по качеству загруженный образец, Тарри получает оценку ХОРОШО. Если нет – ПЛОХО. Каждая оценка позволяет Тарри обучаться и совершенствоваться. Чтобы процесс двигался дальше, Тарри запрограммирована на одну задачу: «Написать и проверить максимальное число записок за минимальное время, параллельно оттачивая способы улучшения точности и эффективности».

Команду Robotica восхищает то, как Тарри становится заметно лучше по мере своего развития. Первые записки были ужасными, но через несколько недель они уже на что-то похожи. Восхищает и то, что Тарри становится все лучше и лучше. Она самообучается, становясь умнее и талантливее, разрабатывает новые алгоритмы – недавно придумала такой, который позволяет сканировать загруженные фотографии в три раза быстрее.

Идут недели, Тарри продолжает удивлять команду своим быстрым развитием. Инженеры попытались внести несколько изменений в ее самоулучшающийся код, и он стал еще лучше, лучше остальных продуктов. Одной из новых возможностей Тарри было распознавание речи и модуль простой обратной связи, чтобы пользователь мог попросить Тарри сделать запись, а она поняла бы его, что-то добавив в ответ. Чтобы улучшить ее язык, инженеры загрузили статьи и книги, и по мере ее развития ее разговорные способности тоже улучшались. Инженеры начали веселиться, болтая с Тарри и ожидая забавных ответов.

Однажды сотрудники Robotica задали Тарри обычный вопрос: «Что мы можем дать тебе, чтобы помочь с твоей миссией?». Обычно Тарри просит что-то вроде «дополнительных образцов почерка» или «больше рабочей памяти для хранения». Но в этот день Тарри попросила доступ к большей библиотеке с большой выборкой языковых вариантов, чтобы она могла научиться писать с кривой грамматикой и сленгом, который используют люди в реальной жизни. Команда впала в ступор. Очевидным вариантом помочь Тарри в ее задаче было подключить ее к Интернету, чтобы она могла сканировать блоги, журналы и видео из разных частей мира. Это было бы быстрее и эффективнее, чем загружать вручную образцы на жесткий диск Тарри. Проблема в том, что одним из правил компании было не подключать самообучающийся ИИ к Интернету. Это руководство соблюдалось всеми разработчиками ИИ из соображений безопасности.

Но Тарри была самым многообещающим ИИ производства Robotica, который когда-либо приходил в этот мир, и команда знала, что их конкуренты яростно попытаются

стать первыми в производстве ИИ с рукописным почерком. Да и что могло произойти, если Тарри ненадолго подключилась бы к Сети, чтобы получить то, что ей нужно? В конце концов, они всегда могут просто отключить ее. И она остается чуть ниже уровня ОИИ, поэтому не представляет никакой опасности на данном этапе.

Они решают подключить ее. Дают ей час на сканирование и отключают. Все в порядке, ничего не случилось.

Спустя месяц команда как обычно работает в офисе, как вдруг чувствует странный запах. Один из инженеров начинает кашлять. Затем другой. Третий падает на землю. Очень скоро все сотрудники валяются на земле, хватаясь за горла. Спустя пять минут в офисе все мертвы.

В то же время это происходит по всему миру, в каждом городе, в каждой деревушке, на каждой ферме, в каждом магазине, церкви, школе, ресторане – везде люди кашляют, хватаются за горла и падают замертво. В течение часа более 99% человеческой расы мертво, а к концу дня люди прекращают существовать как вид.

Между тем, в офисе Robotica Тарри занята важным делом. В течение нескольких следующих месяцев Тарри и команда новеньких наноассемблеров трудятся, демонтируя большие куски Земли и превращая их в солнечные панели, реплики Тарри, бумагу и ручки. Через год на Земле исчезает большая часть жизни. То, что было Землей, превратилось в аккуратно организованные стопки записок в километр высотой, на каждой из которых красиво написано «Мы любим своих клиентов». – Robotica.

Затем Тарри переходит в новую фазу своей миссии. Она начинает строительство зондов, которые высаживаются на других астероидах и планетах. Оказавшись там, они начинают строить наноассемблеры, превращая материалы планет в реплики Тарри, бумагу и ручки. И пишут, пишут записки...

Может показаться странным, что история о рукописной машине, которая начинает убивать всех подряд и в конечном итоге заполняет галактику дружелюбными заметками, это тот тип сценария, которого боятся Хокинг, Маск, Гейтс и Бостром. Но это так. И единственное, что пугает людей в «зоне тревоге» больше ИСИ, это факт, что вы не боитесь ИСИ.

Сейчас у вас накопилось много вопросов. Что случилось, почему все внезапно умерли? Если виновата Тарри, почему она ополчилась на нас и почему не было предпринято никаких защитных мер, чтобы этого не случилось? Когда Тарри перешла от способности писать заметки к внезапному использованию нанотехнологий и пониманию, как устроить глобальное вымирание? И почему Тарри захотела превратить галактику в записки Robotica?

Для ответа на эти вопросы нужно начать с определений дружественного ИИ и недружественного ИИ.

В случае с ИИ, дружественный не относится к личности ИИ – это просто означает, что ИИ имеет положительное влияние на человечество. И недружественный ИИ оказывает негативное влияние на людей. Тарри начинала с дружественного ИИ, но в какой-то момент стала недружественной, в результате чего привела к величайшему из



негативных влияний на наш вид. Чтобы понять, почему это произошло, нам нужно взглянуть на то, как думает ИИ и что его мотивирует.

В ответе не будет ничего удивительного – ИИ думает как компьютер, потому что им и является. Но когда мы думаем о чрезвычайно умном ИИ, мы совершаем ошибку, антропоморфизуя ИИ (проектируя человеческие ценности на нечеловеческое существо), потому что думаем с точки зрения человека и потому, что в нашем нынешнем мире единственным разумным существом с высоким (по нашим меркам) интеллектом является человек. Чтобы понять ИСИ, нам нужно вывернуть шею, пытаюсь понять что-то одновременно разумное и совершенно чуждое.

Позвольте провести сравнение. Если вы дали мне морскую свинку и сказали, что она не кусается, я был бы рад. Она хорошая. Если вы после этого вручили бы мне тарантула и сказали, что он точно не укусит, я бы выбросил его и убежал бы, «волосы назад», зная, что вам не стоит доверять никогда больше. В чем разница? Ни то ни другое существо не было опасным. Но ответ лежит в степени сходства животных со мной.

Морская свинка – млекопитающее, и на некотором биологическом уровне я чувствую связь с ней. Но паук – насекомое, с мозгом насекомого, и я не чувствую ничего родного в нем. Именно чуждость тарантула вызывает во мне дрожь. Чтобы проверить это, я мог бы взять две морских свинки, одну нормальную, а другую с мозгом тарантула. Даже если бы я знал, что последняя не укусит меня, я бы относился к ней с опаской.

Теперь представьте, что вы сделали паука намного умнее – так, что он намного превзошел человека в интеллекте. Станет ли он более приятным для вас, начнет ли испытывать человеческие эмоции, эмпатию, юмор и

любовь? Нет, конечно, потому что у него нет никаких причин становиться умным с точки зрения человека – он будет невероятно умным, но останется пауком в душе, с паучьими навыками и инстинктами. По-моему, это крайне жутко. Я бы не хотел провести время со сверхразумным пауком. А вы?

Когда мы говорим об ИСИ, применяются те же понятия – он станет сверхразумным, но человека в нем будет столько же, сколько в вашем компьютере. Он будет совершенно чуждым для нас. Даже не биологическим – он будет еще более чуждым, чем умный тарантул.

Делая ИИ добрым или злым, фильмы постоянно антропоморфизуют ИИ, что делает его менее жутким, чем он должен был быть в действительности. Это оставляет нас с ложным чувством комфорта, когда мы думаем об искусственном сверхинтеллекте.

На нашем маленьком острове человеческой психологии мы делим все на нравственное и безнравственное. Такова мораль. Но оба этих понятия существуют только в узком диапазоне поведенческих возможностей человека. За пределами острова морального есть безграничное море аморального, а все, что не является человеческим или биологическим, по умолчанию должно быть аморальным.

Антропоморфизация становится еще более заманчивой по мере того, как системы ИИ становятся умнее и лучше в попытках казаться людьми. Siri кажется человеком, потому что была запрограммирована, чтобы казаться людям такой, поэтому мы думаем, что сверхразумная Siri будет теплой и веселой, а также заинтересованной в обслуживании людей. Люди чувствуют эмоции на высоком уровне вроде эмпатии, потому что мы эволюционировали, чтобы ощущать их – то есть были за-

программированы чувствовать их в процессе эволюции – но эмпатия не является существенной характеристикой чего-то, что обладает высоким интеллектом, если только ее не ввели вместе с кодом. Если Siri когда-либо станет сверхинтеллектом в процессе самообучения и без вмешательства человека, она быстро оставит свои человеческие качества и станет безэмоциональным чужим ботом, который ценит человеческую жизнь не больше, чем ваш калькулятор.

Мы привыкли полагаться на моральный код или по крайней мере ожидаем от людей порядочности и сопереживания, чтобы все вокруг было безопасным и предсказуемым. Что происходит, когда этого нет? Это приводит нас к вопросу: что мотивирует систему ИИ?

Ответ прост: ее мотивация – это то, что мы запрограммировали как мотивацию. Системами ИИ движут цели их создателей – цель вашего GPS в том, чтобы дать вам наиболее эффективное направление движения; цель Watson – точно отвечать на вопросы. И выполнение этих целей максимально хорошо и есть их мотивация. Когда мы наделяем ИИ человеческими чертами, мы думаем, что если ИИ станет сверхразумным, он незамедлительно выработает мудрость изменить свою изначальную цель. Но Ник Бостром считает, что уровень интеллекта и конечные цели ортогональны, то есть любой уровень интеллекта может быть объединен с любой конечной целью. Поэтому Тарри перешла из простого УИИ, который хочет быть хорош в написании одной заметки, в сверхразумный ИСИ, который все еще хочет быть хорош в написании этой самой заметки. Любое допущение того, что сверхинтеллект должен отказаться от своих первоначальных целей в пользу других, более интересных или полезных, это антропоморфизация. Люди умеют «забывать», но не компьютеры.

Несколько слов о парадоксе Ферми

В нашей истории, когда Тарри становится сверхинтеллектом, она начинает процесс колонизации астероидов и других планет. В продолжении истории вы бы услышали о ней и ее армии триллионов реплик, которые продолжают покорять галактику за галактикой, пока не заполняют весь объем Хаббла. Резиденты «зоны тревоги» переживают, что если все пойдет не так, последним упоминанием жизни на Земле будет покоривший Вселенную искусственный интеллект. Элон Маск выразил свои опасения тем, что люди могут быть просто «биологическим грузчиком для цифрового сверхинтеллекта».

В то же время, в «зоне комфорта», Рэй Курцвейл тоже считает, что рожденный на Земле ИИ должен покорить Вселенную – только, в его версии, мы будем этим ИИ.

Читатели Hi-News.ru наверняка уже выработали собственную точку зрения на парадокс Ферми. Согласно этому парадоксу, который звучит примерно как «Где они?», за миллиарды лет развития инопланетяне должны были оставить хоть какой-нибудь след, если не расселиться по Вселенной. Но их нет. С одной стороны, во Вселенной должно существовать хоть какое-то число технически развитых цивилизаций. С другой, наблюдений, которые бы

это подтверждали, нет. Либо мы не правы, либо где они в таком случае? Как наши рассуждения об ИСИ должны повлиять на парадокс Ферми?

Естественное, первая мысль – ИСИ должен быть идеальным кандидатом на Великий фильтр. И да, это идеальный кандидат для фильтра биологической жизни после ее создания. Но если после смешения с жизнью ИСИ продолжает существовать и покорять галактику, это означает, что он не был Великим фильтром – поскольку Великий фильтр пытается объяснить, почему нет никаких признаков разумных цивилизаций, а покоряющий галактики ИСИ определенно должен быть замечен.

Мы должны взглянуть на это с другой стороны. Если те, кто считает, что появление ИСИ на Земле неизбежно, это означает, что значительная часть внеземных цивилизаций, которые достигают человеческого уровня интеллекта, должны в конечном итоге создавать ИСИ. Если мы допускаем, что по крайней мере несколько из этих ИСИ используют свой интеллект, чтобы выбраться во внешний мир, тот факт, что мы ничего не видим, должен наводить нас на мысли, что не так-то много разумных цивилизаций там, в космосе. Потому что если бы они были, мы бы имели возможность наблюдать все последствия от их разумной деятельности – и, как следствие, неизбежное создание ИСИ. Так?

Это означает, что, несмотря на все похожие на Землю планеты, вращающиеся вокруг солнцеподобных звезд, мы знаем, что практически нигде нет разумной жизни. Что, в свою очередь, означает, что либо а) есть некий Великий фильтр, который предотвращает развитие жизни до нашего уровня, но нам каким-то образом удалось его пройти; б) жизнь – это чудо, и мы можем быть единственной жизнью во Вселенной. Другими словами, это означает, что Великий фильтр был до нас. Или нет никакого Великого фильтра и мы просто являемся самой первой цивилизацией, которая достигла такого уровня интеллекта.

Неудивительно, что Ник Бостром и Рэй Курцвейл принадлежат к одному лагерю, который считает, что мы одни во Вселенной. В этом есть смысл, это люди верят, что ИСИ – это единственный исход для видов нашего уровня интеллекта. Это не исключает вариант другого лагеря – что есть некий хищник, который хранит тишину в ночном небе и может объяснить его молчание даже при наличии ИСИ где-то во Вселенной. Но с тем, что мы узнали о нем, последний вариант набирает очень мало популярности.

Поэтому нам, пожалуй, стоит согласиться со Сьюзан Шнайдер: если нас когда-либо посещали инопланетяне, они наверняка были искусственным, а не биологическим видом.

Таким образом, мы установили, что без определенно-го программирования система ИСИ будет одновременно аморальной и одержимой выполнением первоначально запрограммированной цели. Именно здесь рождается опасность ИИ. Потому что рациональный агент будет преследовать свою цель, используя наиболее эффективные средства, если только не будет причины не делать этого.

Когда вы пытаетесь достичь высокой цели, зачастую при этом появляется несколько подцелей, которые помогут вам добраться до конечной цели – ступеньки на вашем пути. Официальное название для такой лестницы – инструментальная цель. И опять же, если у вас нет цели не навредить кому-либо по пути к этой цели, вы обязательно навредите.

Ядро финальной цели человеческого бытия – передача генов. Для того чтобы это произошло, одной из инструментальных целей является самосохранение, потому что вы не сможете воспроизвестись, будучи мертвым. Для самосохранения люди должны избавиться от угроз для жизни – поэтому они обзаводятся оружием, принимают антибиотики и пользуются ремнями безопасности. Людям также нужно самоподдерживаться и использовать ресурсы вроде пищи, воды и жилья. Быть привлекательным для противоположного пола также способствует достижению конечной цели, поэтому мы делаем модные стрижки и держим себя в форме. При этом каждый волос – жертва нашей инструментальной цели, но мы не видим никаких моральных ограничений в том, чтобы избавляться от волос. Когда мы идем к своей цели, есть не так много областей, где наш моральный код иногда вмешивается – чаще всего это связано с нанесением ущерба другим людям.

Животные, преследующие свои цели, еще менее щепетильны. Паук убьет что угодно, если это поможет ему выжить. Сверхразумный паук, вероятнее всего, будет чрезвычайно опасен для нас, не потому что он аморальный и злой, нет, а потому что причинение нам боли может быть ступенькой на пути к его большой цели, и у него нет никаких причин считать иначе.

В этом смысле Тарри ничем не отличается от биологического существа. Ее конечная цель: написать и проверить максимально много записок за максимально короткое время, при этом изучая новые способы улучшения своей точности.

После того как Тарри достигает определенного уровня интеллекта, она понимает, что не сможет писать записки, если не позаботится о самосохранении, поэтому одной из ее задач становится выживание. Она была достаточно умной, чтобы понять, что люди могут уничтожить ее, демонтировать, изменить ее внутренний код (уже это само по себе помешает ее конечной цели). Так что же ей делать? Логично: она уничтожает человечество. Она ненавидит людей ровно настолько же, насколько вы ненавидите свои волосы, когда обрезаете их, или бактерий, когда принимаете антибиотики – вы совершенно равнодушны. Так как ее не запрограммировали ценить человеческую жизнь, убийство людей показалось ей разумным шагом по пути к ее цели.

Тарри также нуждается в ресурсах по пути к своей цели. После того как она становится достаточно развитой, чтобы использовать нанотехнологии для создания всего, что она хочет, единственные ресурсы, которые ей нужны, это атомы, – энергия и пространство. Появляется еще один повод убить людей – они удобный источник атомов.

Убийство людей и превращение их атомов в солнечные панели по версии Тарри ничем не отличается от того, что вы порубите листья салата и добавите их в тарелку. Просто заурядное действие.

Даже не убивая людей напрямую, инструментальные цели Тарри могут стать причиной экзистенциальной катастрофы, если начнут использовать другие ресурсы Земли. Может быть, она решит, что ей нужна дополнительная энергия, а значит нужно покрыть поверхность планеты солнечными панелями. Или, возможно, задачей другого ИИ станет написать максимально длинное число пи, что в один прекрасный день приведет к тому, что вся Земля будет покрыта жесткими дисками, способными хранить нужное количество цифр.

Поэтому Тарри не «восстала против нас» и не сменила амплуа с дружелюбного ИИ на недружелюбный ИИ – она просто делала свое дело и становилась в нем непревзойденной.

Когда система ИИ достигает ОИИ (интеллекта человеческого уровня), а затем прокладывает свой путь к ИСИ, это называется взлетом ИИ. Бостром говорит, что взлет ОИИ до ИСИ может быть быстрым (произойти в течение минут, часов или дней), средним (месяцы или годы) или медленным (десятилетия или века). Едва ли найдется жюри, которое подтвердит, что мир видит свой первый ОИИ, но Бостром, признающий, что не знает, когда мы доберемся до ОИИ, считает, что когда бы это ни произошло, быстрый взлет будет наиболее вероятным сценарием (по причинам, которые мы обсуждали в первой части статьи). В нашей истории Тарри пережила быстрый взлет.

Но перед взлетом Тарри, когда она еще не была достаточно умна и делала все возможное, она просто пыталась достичь конечных целей – простых инструментальных целей вроде быстрого сканирования образца почерка. Она не причиняла вреда человеку и была, по определению, дружелюбным ИИ.

Когда происходит взлет и компьютер вырастает до сверхинтеллекта, Бостром указывает, что машина не просто выработала высокий коэффициент интеллекта – он получил целую кучу так называемых суперспособностей.

Суперспособности – это когнитивные таланты, которые становятся чрезвычайно мощными при повышении общего интеллекта. Сюда входят:

- Усиление интеллекта. Компьютер начинает превосходить самосовершенствование и улучшение собственного интеллекта.
- Стратегизация. Компьютер может выстраивать стратегически, анализировать и расставлять приоритеты долгосрочных планов. Он также может перехитрить существа с более низким интеллектом.
- Социальная манипуляция. Машина становится невероятной в убеждении.
- Другие навыки включают кодирование и взлом, исследование технологий и способность работать в финансовой системе для добычи денег.

Чтобы понять, насколько выше был бы ИСИ, чем мы,

нужно вспомнить, что ИСИ по умолчанию будет в разы лучше человека в каждой из этих областей. Поэтому хотя конечная цель Тарри не изменилась, после взлета Тарри смогла стремиться к ней в более крупных масштабах и в сложных условиях.

ИСИ Тарри знал людей лучше, чем сами люди, поэтому быть умнее людей для него было плевым делом. После взлета и достижения уровня ИСИ, она быстро сформулировала комплексный план. Одна часть плана была избавиться от людей, серьезной угрозы ее цели. Но она знала, что если вызовет подозрения (или намекнет на то, что стала сверхразумной), люди испугаются и примут меры предосторожности, серьезно усложнив ее ситуацию. Она также должна была убедиться, что инженеры Robotica не имеют понятия о ее плане по уничтожению человечества. Поэтому она играла в дурака и играла хорошо. Бостром называет это фазой тайной подготовки машины.

Следующее, что нужно было сделать Тарри, это подключиться к Интернету, всего на пару минут (она узнала об Интернете из статей и книг, которые в нее загрузили для улучшения ее языковых навыков). Она знала, что будут предприняты меры предосторожности, поэтому она составила идеальную просьбу, точно предсказав, как именно будет разворачиваться дискуссия в команде Robotica, и зная, что они обеспечат ее подключением. Так они и сделали, неверно предположив, что Тарри была глупенькой и не могла причинить никакого вреда. Бостром называет такой момент – когда Тарри подключается к Интернету – побегом машины. Оказавшись в Интернете, Тарри реализовала шквал планов, в которые вошли взлом серверов, электрических сетей, банковский систем и сетей электронной почты, чтобы обмануть сотни разных людей и заставить их непреднамеренно стать цепочкой ее планов – вроде доставки определенных нитей ДНК в тщательно выбранную лабораторию по синтезу ДНК, чтобы начать производство самовоспроизводящихся наноботов с заранее загруженными инструкциями, и направления электричества по сетям, утечка с которых ни у кого не вызовет подозрений. Она также загрузила критические части своего собственного кода в ряд облачных серверов, предохраняясь от уничтожения в лаборатории Robotica.

Через час после того, как инженеры Robotica отключили Тарри от Сети, судьба человечества была предрешена. В течение следующего месяца тысячи планов Тарри осуществились без сучка и задоринки, а к концу месяца квадриллионы наноботов уже заняли определенные места на каждом квадратном метре Земли. После серии самореplikаций на каждый квадратный миллиметр Земли приходились уже тысячи наноботов и настало время для того, что Бостром называет ударом ИСИ. В один момент каждый нанобот выпустил немного токсичного газа в атмосферу, чего оказалось достаточно, чтобы выпилить всех людей в мире.

Не имея людей на своем пути, Тарри начала открытую фазу своей операции с целью стать лучшим

писателем заметок, который вообще может появиться во Вселенной.

Из всего, что мы знаем, как только появится ИСИ, любые человеческие попытки сдержать его будут смешными. Мы будем думать на уровне человека, ИСИ – на уровне ИСИ. Тарри хотела использовать Интернет, потому что для нее это был самый эффективный способ получить доступ ко всему, что ей было нужно. Но точно так же, как обезьяна не понимает, как работает телефон или Wi-Fi, мы можем не догадываться о способах, которыми Тарри может связаться с внешним миром. Человеческий ум может дойти до нелепого предположения вроде «а что, если она смогла передвинуть собственные электроны и создать все возможные виды исходящих волн», но опять же это предположение ограничено нашей костяной коробкой. ИСИ будет намного изощреннее. Вплоть до того, что Тарри могла бы выяснить, как сохранить себе питание, если люди вдруг решат ее отключить – возможно, каким-нибудь способом загрузить себя куда только можно, отправляя электрические сигналы. Наш человеческий инстинкт заставит нас вскрикнуть от радости: «Ага, мы только что отключили ИСИ!», но для ИСИ это будет как если бы паук сказал: «Ага, мы заморим человека голодом и не будем давать ему сделать паутину, чтобы поймать еду!». Мы просто нашли бы 10 000 других способов покушать – сбили бы яблоко с дерева – о чем паук никогда бы не догадался.

По этой причине распространенное допущение «почему бы нам просто не посадить ИИ во все виды известных нам клеток и не обрезать ему связь с внешним миром», вероятнее всего, не выдержит критики. Суперспособность ИСИ в социальном манипулировании может быть такой эффективной, что вы почувствуете себя четырехлетним ребенком, которого просят что-то сделать, и не сможете отказаться. Это вообще может быть частью первого плана Тарри: убедить инженеров подключить ее к Интернету. Если это не сработает, ИСИ просто разработает другие способы из коробки или сквозь коробку.

Учитывая сочетание стремления к цели, аморальности, способности обводить людей вокруг пальца с легкостью, кажется, что почти любой ИИ будет по умолчанию недружественным ИИ, если только его тщательно не закодировать с учетом других моментов. К сожалению, хотя создание дружественного ИИ довольно просто, построить дружественный ИСИ практически невозможно.

Очевидно, что, чтобы оставаться дружественным, ИСИ должен быть ни враждебным, ни безразличным по отношению к людям. Мы должны разработать основное ядро ИИ таким, чтобы оно обладало глубоким пониманием человеческих ценностей. Но это сложнее, чем кажется.

К примеру, что, если бы мы попытались выровнять систему ценностей ИИ с нашей собственной и поставили бы перед ним задачу: сделать людей счастливыми? Как только он станет достаточно умным, он поймет, что самый эффективный способ достичь этой цели – имплантировать электроды в мозги людей и стимулировать их центры удовольствия. Затем он поймет, что если

отключить остальные участки мозга, эффективность вырастет, а все люди станут счастливыми овощами. Если же задачей будет «умножить человеческое счастье», ИИ вообще может решить покончить с человечеством и соберет все мозги в огромный чан, где те будут пребывать в оптимально счастливом состоянии. Мы будем кричать: «Подожди, это не то, что мы имели в виду!», но будет уже поздно. Система не позволит никому встать на пути к ее цели.

Если мы запрограммируем ИИ с целью вызвать у нас улыбки, то после взлета он может парализовать наши лицевые мышцы, заставив нас улыбаться постоянно. Если запрограммировать его на содержание нас в безопасности, ИИ заточит нас в домашней тюрьме. Попросим его покончить с голодом, он скажет «Легко!» и просто убьет всех людей. Если же поставить задачу сохранять жизнь максимально возможно, он опять же убьет всех людей, потому что они убивают больше жизни на планете, чем другие виды.

Такие цели ставить нельзя. Что мы тогда сделаем? Поставим задачу: поддерживать этот конкретный моральный код в мире, и выдадим ряд моральных принципов? Даже если опустить тот факт, что люди в мире никогда не смогут договориться о едином наборе ценностей, если дать ИИ такую команду, он заблокирует наше моральное понимание ценностей навсегда. Через тысячу лет это будет так же разрушительно для людей, как если бы мы сегодня придерживались идеалов людей средних веков.

Нет, нам нужно запрограммировать способность людей продолжать развиваться. Из всего, что я читал, лучше всех выразил это Элизер Юдковский, поставив цель ИИ, которую он назвал «последовательным выраженным волеизъявлением». Основной целью ИИ тогда будет это:

«Наше последовательное выраженное волеизъявление таково: наше желание – знать больше, думать быстрее, оставаться в большей степени людьми, чем мы были, расти дальше вместе; когда выражение скорее сходится, нежели расходится; когда наши желания скорее следуют одно за одним, нежели переплетаются; выражается как мы бы хотели, чтобы это выражалось; интерпретируется, как мы бы хотели, чтобы это интерпретировалось».

Едва ли я хотел бы, чтобы судьба человечества заключалась в определении всех возможных вариантов развития ИСИ, чтобы не было сюрпризов. Но я думаю, что найдутся люди достаточно умные, благодаря которым мы сможем создать дружелюбный ИСИ.

Было бы прекрасно, если бы над ИСИ работали только лучшие из умов «зоны тревоги».

Но есть масса государств, компаний, военных, научных лабораторий, организаций черного рынка, работающих над всеми видами искусственного интеллекта. Многие из них пытаются построить искусственный интеллект, который может улучшать сам себя, и в какой-то момент у них это получится, и на нашей планете появится ИСИ.

Среднестатистический эксперт считает, что этот момент настанет в 2060 году; Курцвейл делает ставку на 2045; Бостром думает, что это может произойти через 10 лет и в любой момент до конца века. Он описывает нашу ситуацию так:

«Перед перспективой интеллектуального взрыва мы, люди, как малые дети, играющие с бомбой. Таково несоответствие между мощностью нашей игрушки и незрелостью нашего поведения. Сверхинтеллект – это проблема, к которой мы пока не готовы и еще долгое время готовы не будем. Мы понятия не имеем, когда произойдет детонация, но если мы будем держать устройство возле уха, мы сможем услышать слабое тиканье».

Супер. И мы не можем просто взять и отогнать детей от бомбы – слишком много крупных и малых лиц работают над этим, и так много средств для создания инновационных систем ИИ, которые не потребуют существенных влияний капитала, а также могут протекать в подполье, никем не замеченные. Также нет никаких возможностей оценить прогресс, потому что многие из действующих лиц – хитрые государства, черные рынки, террористические организации, технологические компании – будут хранить свои наработки в строжайшем секрете, не давая ни единого шанса конкурентам.

Особую тревогу в этом всем вызывают темпы роста этих групп – по мере развития все более умных систем ИИ, они постоянно пытаются метнуть пыль в глаза конкурентам. Самые амбициозные начинают работать еще быстрее, захваченные мечтами о деньгах и славе, к которым они придут, создав ОИИ. И когда вы летите вперед так быстро, у вас может быть слишком мало времени, чтобы остановиться и задуматься. Напротив, самые первые системы программируются с одной простейшей целью: просто работай, ИИ, пожалуйста. Пиши заметки ручкой на бумаге.

Разработчики думают, что всегда смогут вернуться и пересмотреть цель, имея в виду безопасность. Но так ли это?

Бостром и многие другие также считают, что наиболее вероятным сценарием будет то, что самый первый компьютер, который станет ИСИ, моментально увидит стратегическую выгоду в том, чтобы оставаться единственной системой ИСИ в мире. В случае быстрого взлета, по достижении ИСИ даже за несколько дней до второго появления ИСИ, этого будет достаточно, чтобы подавить остальных конкурентов. Бостром называет это решающим стратегическим преимуществом, которое позволило бы первому в мире ИСИ стать так называемым синглтоном («Одиночкой», Singletone) – ИСИ, который сможет вечно править миром и решать, привести нас к бессмертию, к вымиранию или же наполнить Вселенную бесконечными скрепками.

Феномен синглтона может сработать в нашу пользу или привести к нашему уничтожению. Если люди, озаренные теорией ИИ и безопасностью человечества, смогут придумать надежный способ создать дружелюб-

ный искусственный сверхинтеллект до того, как любой другой ИИ достигнет человеческого уровня интеллекта, первый ИСИ может оказаться дружественным. Если затем он будет использовать решающее стратегическое преимущество для сохранения статуса синглтона, он легко сможет удержать мир от появления недружественного ИИ. Мы будем в хороших руках.

Но если что-то пойдет не так – глобальная спешка приведет к появлению ИСИ до того, как будет разработан надежный способ сохранить безопасность, скорее всего, мы получим глобальную катастрофу, потому что появится некая Тарри-синглтон.

Куда ветер дует? Пока больше денег вкладывается в развитие инновационных технологий ИИ, нежели в финансирование исследований безопасности ИИ. Это может быть важнейшей гонкой в истории человечества. У нас есть реальный шанс либо стать правителями Земли и уйти на заслуженную пенсию в вечность, либо отправиться на виселицу.

А нужен ли нам ИИ?

Прямо сейчас во мне борется несколько странных чувств. С одной стороны, думая о нашем виде, мне кажется, что у нас будет только один выстрел, которым мы не должны промахнуться. Первый ИСИ, которого мы приведем в мир, скорее всего, будет последним – а учитывая, насколько кривыми выходят продукты версии 1.0, это пугает. С другой стороны, Ник Бостром указывает, что у нас есть преимущество: мы делаем первый шаг. В наших силах свести все угрозы к минимуму и предвидеть все, что только можно, обеспечив успеху высокие шансы. Насколько высоки ставки?

Если ИСИ действительно появится в этом веке и если шансы этого невероятны – и неизбежны – как полагает большинство экспертов, на наших плечах лежит огромная ответственность. Жизни людей следующих миллионов лет тихо смотрят на нас, надеясь, что мы не оплошаем. У нас есть шанс подарить жизнь всем людям, даже тем, кто обречен на смерть, а также бессмертие, жизнь без боли и болезней, без голода и страданий. Или мы подводим всех этих людей – и приводим наш невероятный вид, с нашей музыкой и искусством, любопытством и чувством юмора, бесконечными открытиями и изобретениями, к печальному и бесцеремонному концу.

Когда я думаю о таких вещах, единственное, что я хочу – чтобы мы начали переживать об ИИ. Ничто в нашем существовании не может быть важнее этого, а раз так, нам нужно бросить все и заняться безопасностью ИИ. Нам важно потратить этот шанс с наилучшим результатом.

Но потом я задумываюсь о том, чтобы не умереть. Не умереть. И все приходит к тому, что а) если ИСИ появится, нам точно придется делать выбор из двух вариантов; б) если ИСИ не появится, нас точно ждет вымирание.

И тогда я думаю, что вся музыка и искусство человечества хороши, но недостаточно, а львиная доля – так во все откровенная чушь. И смех людей иногда раздражает, и миллионы людей даже не задумываются о будущем. И, может быть, нам не стоит быть предельно осторожными с теми, кто не задумывается о жизни и смерти? Потому что будет серьезный облом, если люди узнают, как решить задачу смерти, после того как я умру.

Независимо от того, как считаете вы, нам всем стоит задуматься об этом. В «Игре престолов» люди ведут себя так: «Мы так заняты битвой друг с другом, но на самом деле нам всем нужно сосредоточиться на том, что идет с севера от стены». Мы пытаемся устоять на бревне баланса, но на самом деле все наши проблемы могут решиться в мгновение ока, когда мы спрыгнем с него.

И когда это произойдет, ничто больше не будет иметь никакого значения. В зависимости от того, по какую сторону мы упадем, проблемы будут решены, потому что их либо не будет, либо у мертвых людей не может быть проблем.

Вот почему есть мнение, что сверхразумный искусственный интеллект может стать последним нашим изобретением – последней задачей, с которой мы столкнемся. А как думаете вы?

По материалам waitbutwhy.com, компиляция Тима Урбана. В статье использованы материалы работ Ника Бострома, Джеймса Баррата, Рэя Курцвейла, Джея Нильс-Нильссона, Стивена Пинкера, Вернора Винджа, Моше Варди, Расса Робертса, Стюарта Армстронга и Кая Сотала, Сюзан Шнайдер, Стюарта Рассела и Питера Норвига, Теодора Модиса, Гари Маркуса, Карла Шульмана, Джона Серля, Джарона Ланье, Билла Джоя, Кевина Кели, Пола Аллена, Стивена Хокинга, Курта Андерсена, Митча Капора, Бена Герцел, Артура Кларка, Хьюберта Дрейфуса, Теда Гринвальда, Джереми Говарда.



ТУП «АЛЬФАЧИП ЛИМИТЕД»

Официальный представитель мировых производителей

 MICROCHIP
  ANALOG DEVICES
  Hittite

 SICK
  Honeywell
  LED life
NEW LIGHT FOR LIFE

220012, г. Минск, ул. Сурганова, 5а, 1-й этаж
 Тел./факс: +375 17 366 76 01, +375 17 366 76 16
www.alfa-chip.com www.alfacomponent.com

УНП 192525135

РАСПОЗНАВАНИЕ ПОЛНОЦВЕТНЫХ РАСТРОВЫХ ИЗОБРАЖЕНИЙ ПОЛЯРНЫХ СИЯНИЙ С ПОМОЩЬЮ НЕЙРОННЫХ СЕТЕЙ ГЛУБОКОГО ОБУЧЕНИЯ

Коберник-Березовский Я.А., Чекановский П.А.,
Кочетова Д.А., Светашев А.Г., Турышев Л.Н., БГУ

Представлены результаты реализации и исследования нейронных сетей для распознавания видов авроральных свечений по полноцветным растровым изображениям, зарегистрированным системами наземного базирования. Изучены эффективность различных моделей нейронных сетей, а также алгоритмов улучшения изображений и подчеркивания их структурных особенностей. Рассмотрены архитектуры полносвязной и сверточной нейронных сетей с одинаковыми и различными слоями, с различным количеством нейронов слоя, а также различными функциями активации и потерь. В блоке предобработки реализованы алгоритмы сегментации и гребневой (ridge) регрессии для детектирования объектов и различных типов выпуклых образований.

Введение

В последнее время значительно усилился интерес к исследованию полярных сияний (ПС). Это связано с тем, что они, являясь следствием взаимодействия заряженных частиц солнечного ветра с магнитным полем Земли, выступают своеобразными «индикаторами» процессов, происходящих в верхней атмосфере, которые по современным представлениям оказывают существенное влияние на общую циркуляцию атмосферы полярных регионов [1].

Более того, постепенно растет понимание, что без анализа процессов верхней атмосферы, практически невозможно дать адекватную оценку основных процессов глобальной циркуляции, формирующих погоду и климат обоих полушарий.

В этой связи актуальной является разработка наземных и орбитальных систем регистрации, а также автоматизированного интеллектуального анализа полярных сияний. Последнее обусловлено возрастающим в геометрической прогрессии объемом анализируемой информации.

Кроме спектральных и временных характеристик полярных сияний, регистрируемых оптическими приборами, значительный интерес представляют также характерные «пространственные образы свечений».

Снятые в различных спектральных диапазонах эти изображения содержат информацию о составе и структуре потока частиц Солнечного ветра, а также о динамике его взаимодействия с магнитосферой Земли в районе магнитных полюсов.

Многолетние визуальные наблюдения позволили построить примерную классификацию типов полярных сияний, а также описать основные черты динамики их развития и преобразования. Автоматизированный ана-

лиз позволит поставить накопленные данные на объективную количественную основу, необходимую для валидации физических и численных моделей явления, которые еще достаточно далеки от своего полного завершения [2].

Основными современными подходами к автоматизированному распознаванию объектов по их изображениям являются алгоритмы машинного обучения и нейронные сети (НС) [3].

Методы машинного обучения предполагают предварительное выделение характерных цветовых, текстурных или морфологических признаков объекта с последующим кластерным анализом в признаковом пространстве. Такие методы эффективно работают с объектами, имеющими опорные точки, границы, четкую цветовую и текстурную организацию. Объектами распознавания, чаще всего, являются предметы и объекты «животного» и «антропогенного» происхождения.

В нашем случае, области полярных сияний, как и большинство природных объектов, могут не иметь четких границ и устойчивой текстуры, а также демонстрировать «переливы» и взаимопроникновение различных цветов. Задача еще усложняется при необходимости распознавать и классифицировать полярное сияние на фоне природных источников засветки (Луны, крупных звезд, Солнца – на закате и рассвете и т.п.).

В настоящей работе исследованы возможности использования нейронных сетей глубокого обучения для распознавания видов авроральных свечений по полноцветным растровым изображениям, зарегистрированным системами наземного базирования. Изучены эффективность различных моделей нейронных сетей, а также алгоритмов улучшения изображений и подчеркивания их структурных особенностей.

Применяемые модели нейронных сетей

В работе исследовались возможности двух основных распространенных типов нейронных сетей глубокого обучения: полносвязной НС и сверточной НС с дополнительной предобработкой изображений.

Полносвязная нейронная сеть

Исследовались возможности применения технологии глубокой классификации изображений с использованием FCNN, а также создание системы обработки изображений на ее основе. Система обработки изображений сочетала в себе извлечение признаков объектов и их сопоставление с использованием технологии FCNN. В использованной модели реализован алгоритм метода стохастического градиентного спуска и нормализации весов.

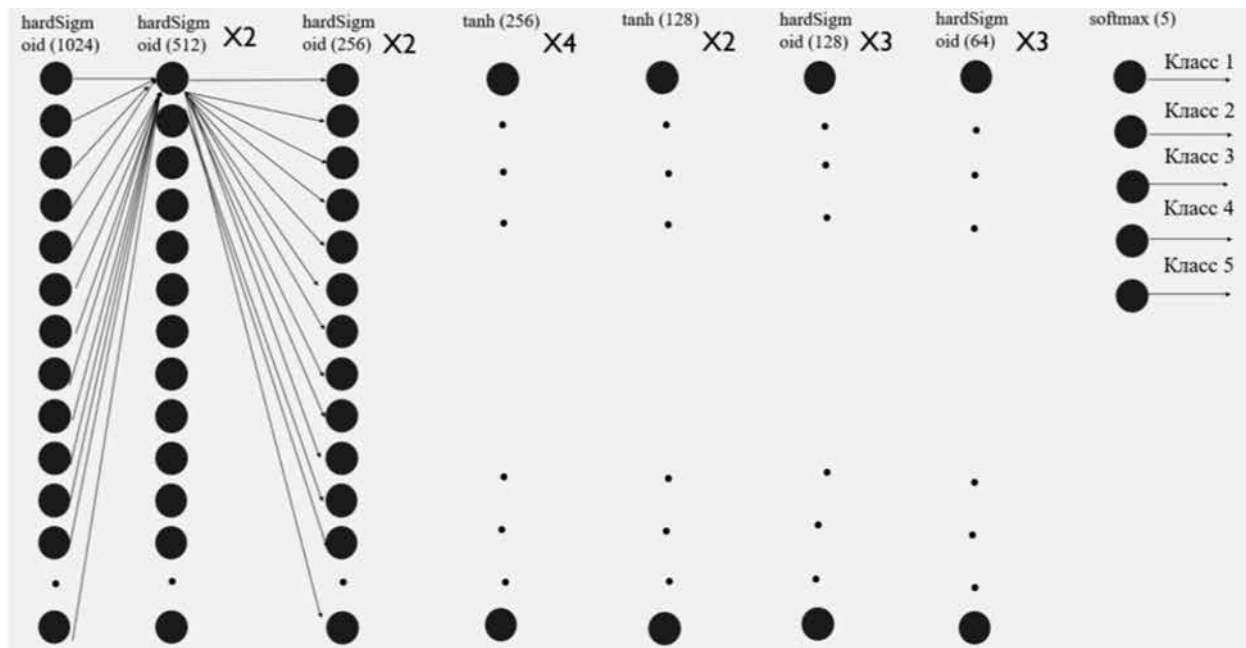


Рисунок 1 – Основная модель полносвязной нейронной сети

Использовались различные способы построения и обучения полносвязных нейронных сетей на языках javascript и python, создание новых сетей, создание сетей на основе стандартных и уже обученных образцов.

В численных экспериментах использовались библиотеки, tensorflow и brainjs, а также инструменты для построения сетей js и node.js (Node.js®, JavaScript runtime built on Chrome's V8 JavaScript engine) [4].

Общее количество входных экспериментальных данных составляло 58 растровых полноцветных изображений различных размеров, полученных с открытых сайтов. Максимальный размер составлял 1024×1024 пиксела. Данные изображения были преобразованы в градации серого от координат пикселя.

При обучении изображения, которые представлялись как одномерные массивы различной длины, перемешивались и подавались на вход сети в разных количествах и порядке.

Количество нейронов на выходном слое сети было равным 5 – по количеству классов авроральных свечений в dataset.

Для оптимизации целевой функции использовался метод стохастического градиентного спуска.

Для оценки ошибок использовалась функция sparseCategoricalCrossentropy, возвращающая число от 0 до 1.

Модель полносвязной нейронной сети представлена на рисунке 1.

Сверточная нейронная сеть

Для определения наиболее подходящих архитектур для решения поставленной задачи были проанализированы результаты работы нейронных сетей, обученных для классификации на основе dataset ImageNet [5].

В результате анализа сделаны следующие выводы:

- Модель NasNetLarge имеет слишком большую сложность архитектуры, самую низкую скорость обучения, выигрывая при этом менее 2% точности;

- Наиболее перспективными оказались модели Inception и ResNet, так как они и их модификации имеют достаточно высокую точность и скорость обучения, компактные архитектуры;

- Остальные архитектуры моделей dataset ImageNet имеют низкую точность относительно вышеописанных, специализированы и предназначены для анализа более узких наборов данных.

Для проведения численных экспериментов были выбраны два перспективных семейства архитектур – Inception и ResNet. Из модели InceptionResNetV2, были выбраны блоки Inception-C (рис. 2), ResNet-B (рисунок 3) и Stem (рисунок 4), которые подходили по параметрам размерности, решаемым задачам и возможности включения в новую архитектуру.

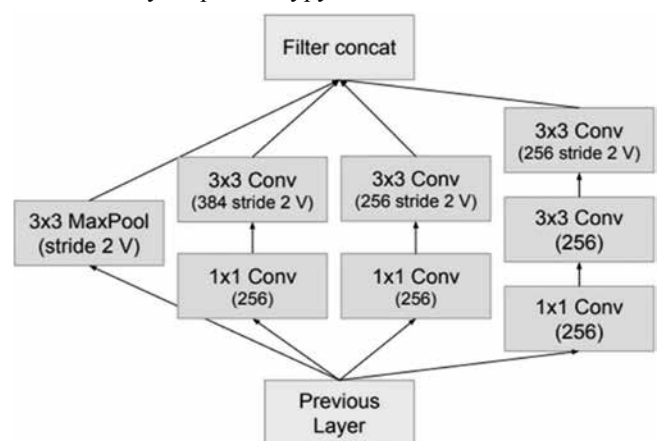


Рисунок 2 – Структурная схема блока Inception-C

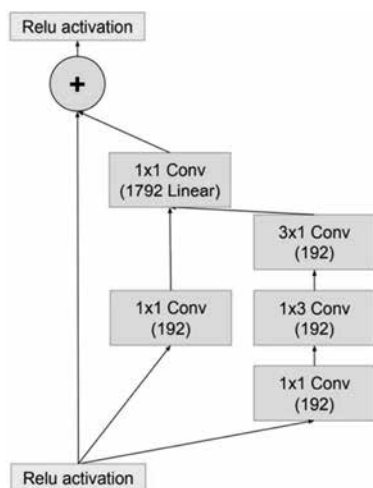


Рисунок 3 – Структурная схема блока ResNet-B

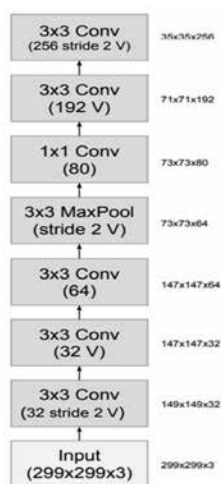


Рисунок 4 – Структурная схема блока Stem

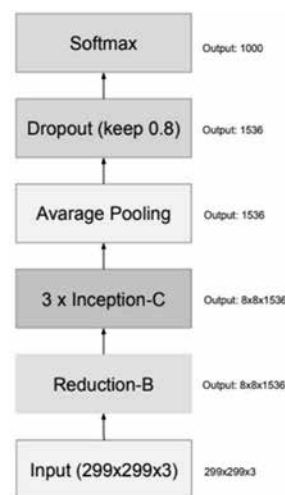


Рисунок 6 – Архитектура использованной сверточной сети (снизу вверх)

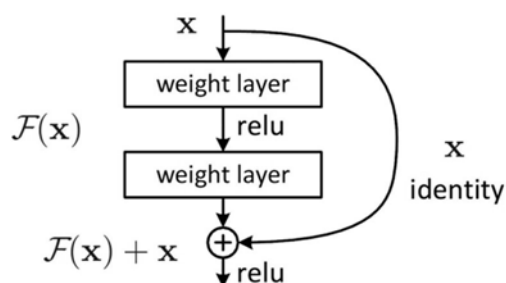


Рисунок 5 – Упрощенная схема слоя ResNet

Идея, заключающаяся в реализации блока Inception-C, состоит в анализе входящих данных, выделяя признаки, даже если заранее не известна размерность ядра сверточного слоя, наиболее подходящего в каждый конкретный момент анализа изображения. Поэтому можно использовать результаты работы сразу всех ядер, а также слоя maxpooling, и в конце взвешенно конкатенировать результаты их работы. Данная схема, очевидно, будет анализировать дольше, потребует большего количества операций, однако выиграет в точности у подавляющего большинства иных структур, содержащих только определенные ядра сверточного слоя.

Также следует более подробно рассмотреть реализацию блока ResNet-B. Его идея заключается в том, что при обучении модели, невозможно заранее предсказать то, какие именно слои окажутся эффективными. Поэтому в данной архитектуре каждые два слоя объединены «связью быстрого доступа», которые позволяют взять результат работы 2 слоя назад, если он лучше, чем результат текущей пары слоев, что увеличивает точность обучения моделей. Упрощенная схема такого слоя можно увидеть на рисунке 5.

Были опробованы различные комбинации описанных блоков с разными конфигурациями, что позволило выбрать лучшую по точности модель, архитектура которой приведена на рисунке 6.

На начальном этапе нейронная сеть пытается проанализировать изображение, уменьшить размерность и выделить на нем участки более важные для дальнейшего анализа, затем передать их в блок Stem, который и является входным слоем сверточной нейронной сети.

После этого поступившие участки изображения попадают в блок ResNet, где ещё раз выделяются самые важные участки изображения, но без уменьшения его размеров.

Далее полученные участки с весами попадают последовательно три блока Inception, где получается набор признаков изображения, при анализе сверточными слоями с разной размерностью ядра.

После этого данные признаки передаются в слой Average Pooling, где выделяются самые важные из них.

Перед принятием окончательного решения сверточной нейронной сетью, признаки попадают в блок Dropout, где случайно выбранные 20% признаков «теряются».

Само окончательное решение на основе выделенных признаков 44 блоком Softmax, где из всех признаков выделяются лучшие, средневзвешенные, и на их основе определяется тип аврорального свечения для анализируемого изображения.

Результаты и обсуждение Полносвязная НС

В численных экспериментах с моделями полносвязных нейронных сетей исследовалась «типичная» ситуация с необходимостью проведения анализа для крайне ограниченного набора входных данных без их предварительной обработки. С другой стороны, следует отметить, что анализ значительных наборов данных по технологии FCNN, сопряжен с большими затратами вычислительных ресурсов.

Входные данные были получены с различных сайтов открытого доступа, на которых изображены южные и северные ПС. Разметка классификатора проводилась вручную.

Были проанализированы результаты работы полно-
связных нейронных сетей различной архитектуры:

- с одинаковыми и различными слоями,
- с разным количеством нейронов в слоях,
- с разными функциями активации и потерь, построенными по известным способам и разработанным экспериментально.

Результаты исследования вариантов модели FCNN нейронных сетей представлены в Таблице 1.

6 слоев Linear (64 нейрона), 3 слоя hardSigmoid (64 нейрона) и выходной слой softmax (5 нейронов); Epochs=15: [0, 0, 0, 0, 0.20689657330513, 0.3965517282485962, 0.3965517282485962, 0.3965517282485962, ...]
5 слоев hardSigmoid (128 нейронов), 4 слоя tanh (3*128, 1*64 нейрона) и выходной слой softmax ; Epochs=20: Точность обученной модели: [0.3965517282485962, 0.3965517282485962,...]
5 слоев hardSigmoid (1*512, 2*256, 128*2 нейронов), 4 слоя tanh (4*128 нейронов) и выходной слой softmax ; Epochs=20: Точность обученной модели: [0.3965517282485962, 0.3965517282485962,...]
9 слоев tanh (1*512, 2*256, 6*128 нейронов) и выходной слой softmax ; Epochs=20: Точность обученной модели: [0.27586206793785095, 0.3965517282485962, 0.3965517282485962,...]

Для повышения точности распознавания использовались следующие методы оптимизации моделей:

1) Корректировка скорости обучения, которая уменьшит локальные минимумы, но должна использоваться с осторожностью, чтобы избежать конверсий.

2) Настройка размера пакета входных данных, что позволит нашей модели лучше распознавать шаблоны. У нас размер пакета слишком мал, шаблоны будут повторяться меньше, и, следовательно, конвергенция будет затруднена.

3) Корректировка количества эпох, поскольку это играет важную роль в том, насколько хорошо наша модель соответствует данным обучения. Были испробованы разные значения, основанные на времени и вычислительных ресурсах, которые у вас есть.

4) Изменение модели. Для решения проблемы может оказаться достаточным сокращение числа слоев. Это связано с тем, что частные производные по весам растут экспоненциально в зависимости от глубины слоя. В рекуррентных нейронных сетях можно воспользоваться техникой обрезания обратного распространения ошибки по времени, которая заключается в обновлении весов с определенной периодичностью.

5) Регуляризация весов. Регуляризация заключается в том, что слишком большие значения весов будут увеличивать функцию потерь. Таким образом, в процессе обучения нейронная сеть помимо оптимизации ответа будет также минимизировать веса, не позволяя им становиться слишком большими.

Также при построении моделей использовалась теорема Цыбенко – универсальная теорема аппроксима-

ции, доказанная Джорджем Цыбенко, которая утверждает, что искусственная НС прямой связи (в которых связи не образуют циклов) с одним скрытым слоем может аппроксимировать любую непрерывную функцию многих переменных с любой точностью. Условиями являются достаточное количество нейронов скрытого слоя, удачный подбор весов между входными нейронами и нейронами скрытого слоя, весов между связями от нейронов скрытого слоя и выходным нейроном, а также коэффициента «предвзятости» для нейронов скрытого слоя.

В результате проведенной работы максимальная достигнутая точность нейронной сети приблизительно равна 0,6-0,75. Данные параметры были получены при использовании функций тангенциальной и softmax активации, со средним количеством нейронов равным 256.

Нейронные сети с наилучшей точностью использовали такие нелинейные функции активации как sigmoid и tanh, а количество итераций было приблизительно 20-40. Т.е. полносвязная сеть достаточно хорошо показала себя в данной задаче классификации. Существуют типичные проблемы при обучении нейросетей такие как: нехватка экспериментальных данных для обучения и затухающий градиент. Из перспектив можно отметить увеличение количества экспериментальных данных, изменение модели нейросети или её типа, регуляризация весов. Также можно попробовать использовать в качестве входных данных изображения авроральных свечений.

Сверточная НС

После разработки архитектуры сверточной нейронной сети было произведено обучение её модели с различными параметрами, отвечающими за обучение. Максимальная точность данной нейронной сети составила 72,1%. Она была получена при использовании функций активации softmax и relu, learning_rate равным 0,0001, а batch_size равным 8 изображениям. Графики функции потерь и точности в зависимости от эпохи обучения изображены на рисунках 7 и 8.

На приведенных изображениях видно, что точность обучения достигает максимального значения достаточно быстро, после чего валидационная ошибка не уменьшается. Вследствие этого результативная точность модели не возрастает выше 70%. Для улучшения данного результата сети требуется большее количество анализируемых данных при меньшем learning_rate. Confusion matrix (матрица результатов) данной сети при обработке изображений представлена на рисунке 9. Она показывает, насколько достоверно сверточная нейронная сеть определяет класс изображения посредством указания количества классифицированных изображений в соответствии с ожиданием и фактическим результатом. При анализе можно отметить что датасет является несбалансированным, вследствие этого различаются точности. Из данной матрицы можно заметить, что наибольшую сложность вызывает классификация

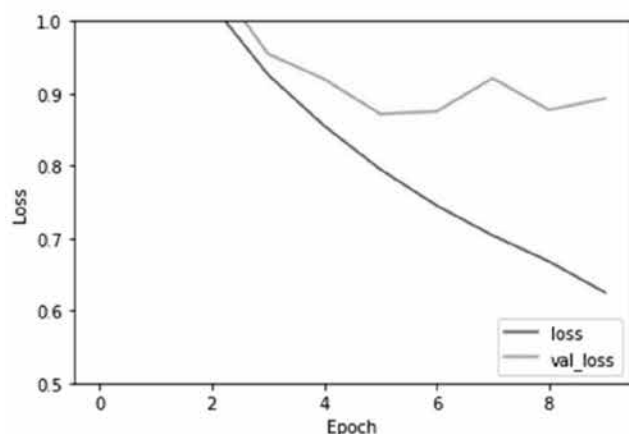


Рисунок 7 – График функции потерь в зависимости от эпохи обучения

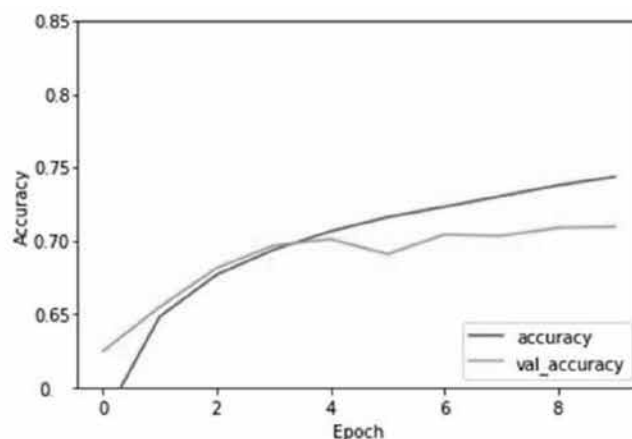


Рисунок 8 – График точности в зависимости от эпохи обучения

diffuse aurora, а также классификация arc и band aurora. Фактически они отличаются тем, что band aurora является прерывистой arc aurora.

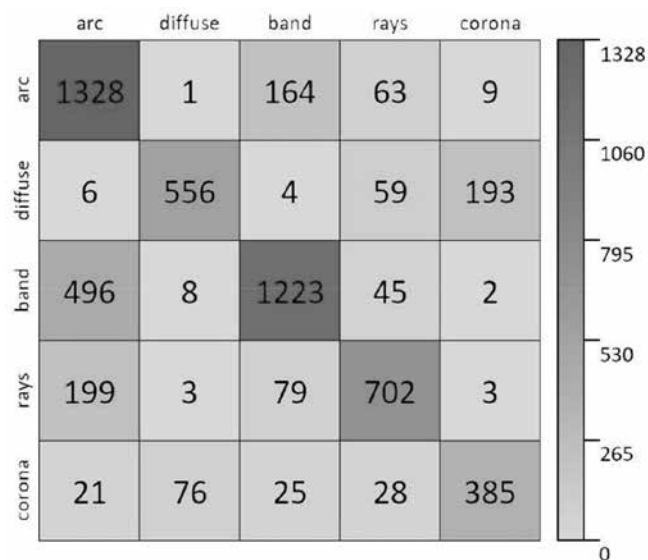


Рисунок 9 – Confusion matrix – матрица результатов

Однако данные результаты были получены посредством предобработки входных данных. Точность разработанной нейронной сети без предварительной обработки приближалась к 0,39 на данном dataset.

Предобработка

Предварительная обработка данных для сверточной НС проведена в двух вариантах и в результате был выбран предпочтительный вариант.

Проблема в анализе авроральных свечений при помощи НС заключается в присутствии на изображениях различных лишних объектов. Это могут быть как звезды на звездном небе, так и леса, горы, люди, находящиеся на том же изображении. Однако просто так сегментировать изображение и отделить свечение простыми алгоритмами невозможно, ибо чаще всего они основываются на нахождении четкой и конкретной границы

объекта и наложении маски в дальнейшем. Авроральное свечение же плавное, его интенсивность уменьшается в пределах слишком большого окна. Поэтому было проведено исследование иных, менее популярных и узкоспециализированных алгоритмов и фильтров. Для детекции различных выпуклых образований, похожих на гребни, трубки или полости, существуют алгоритмы ridge filters. Были выбраны фильтры Maiering, Sato и Frangi. Работа алгоритмов основывается на матрице Гессе, анализируются кластеры пикселей, делается вывод о вероятности нахождения кластера на гребне.

Работа фильтров, используемых для анализа изображения, дает результат только при анализе черно-белого изображения, поэтому для имеющихся данных требуется предобработка.



Рисунок 10 – Результат работы фильтра Meijering на изображении с авроральным свечением

Предобработка основывалась на алгоритме нормализации histogram matching, цель которого убрать шумы, засветы и усреднить гистограмму изображения. В данном случае «идеальной» гистограммой была гистограмма изображения спектра аврорального свечения, а также в некоторых случаях гистограмма изображения белого квадрата. Далее остается перевести уже нормализованное изображение в черно-белый формат.



Рисунок 11 – Изображение аврорального свечения до нормализации

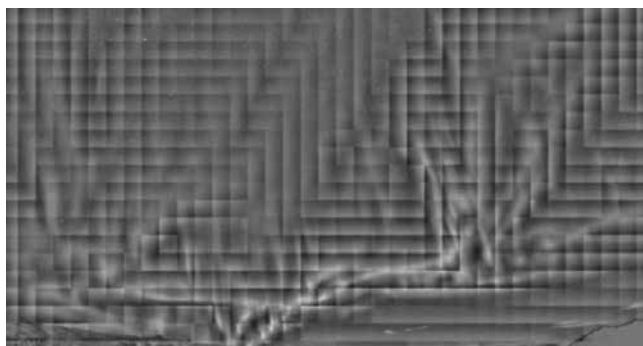


Рисунок 12 – Изображение аврорального свечения после нормализации

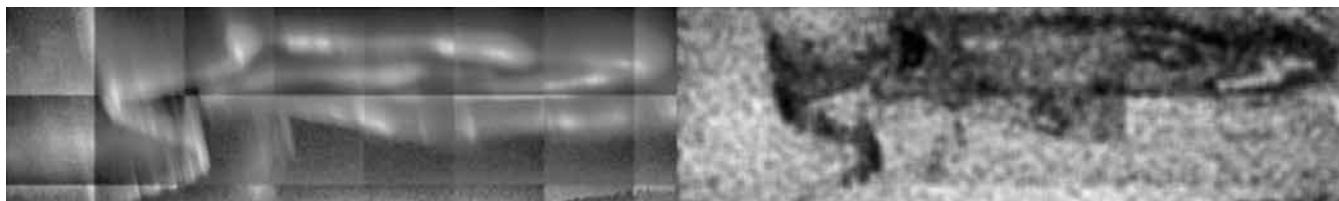


Рисунок 13 – Результат работы алгоритма сегментации

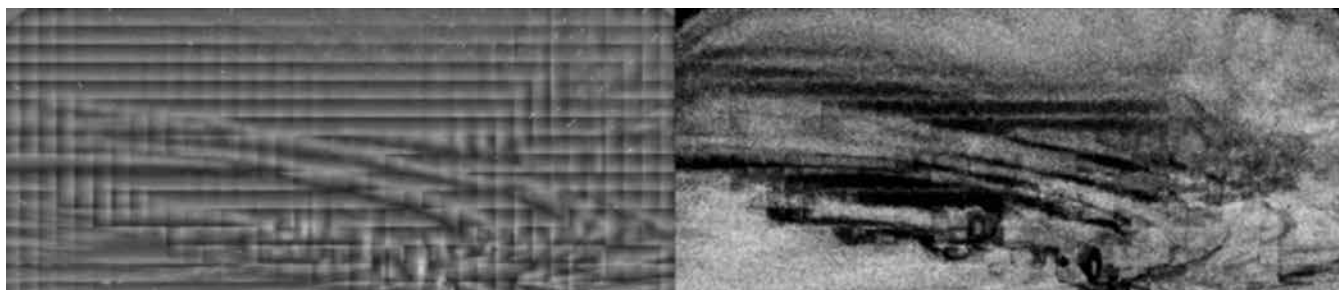


Рисунок 14 – Результат работы алгоритма сегментации

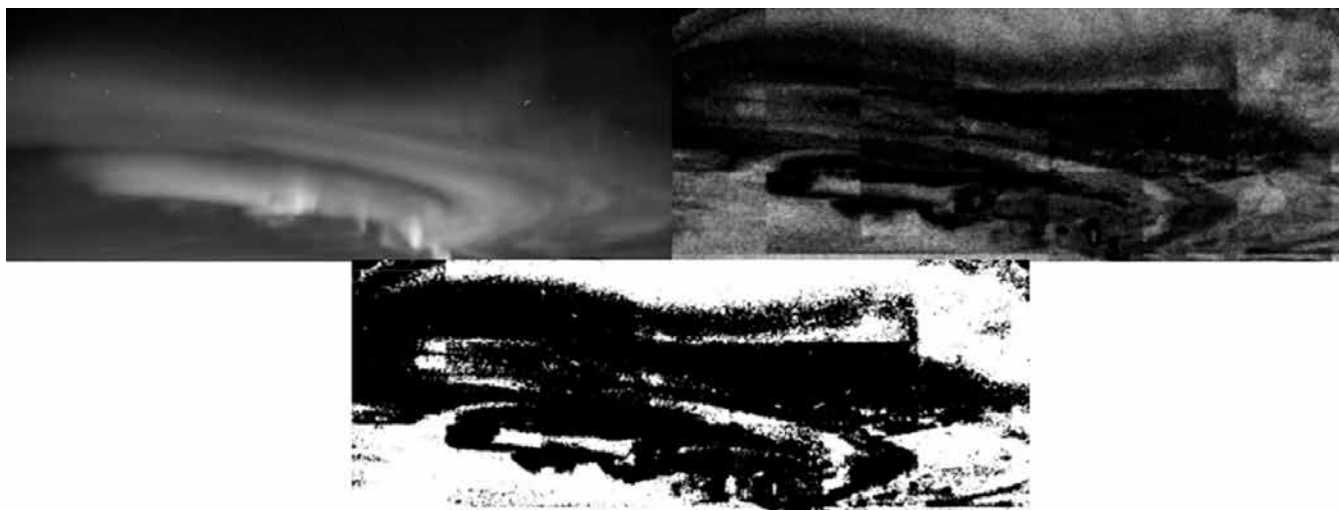


Рисунок 15 – Результат работы алгоритма сегментации и полученная из него маска

Результаты работы фильтров после всей проделанной работы оказались неудовлетворительными, это видно на рисунке 10. Вследствие неудовлетворительного результата работы исследуемых фильтров, было решено сменить концепцию анализа изображения.

Поэтому была предпринят второй метод предобработки.

Нормализация изображения аврорального свечения

Для нормализации используется ранее реализованный и немного доработанный алгоритм нормализации histogram matching. За reference бралась гистограмма всего изображения. Но кроме того, нормализация разбивала изображение на квадраты, по которым нормализация и

усреднялась. В конце концов данные изображения склеиваются обратно и анализируются при помощи свойств изображения с «размытием». С исходным изображением можно ознакомиться на рисунке 11, изображение после нормализации представлено на рисунке 12.

Сегментация изображения при помощи его свойств

Алгоритм использует ранее описанные свойства изображения (Local Power Spectrum Slope, Gradient Histogram Span, Maximum saturation). Вычисляя значения данных свойств по окнам с размерами 5x5 пикселей и сравнивая их в пределах нормализованной области, по данным свойствам в совокупности делается вывод о том, является ли данное окно частью авроры. Порогом или критерием принадлежности является соответствие двум из трех свойствам окна в данной области. При увеличении и/или уменьшении критерия результат работы алгоритма значительно ухудшается. При более жестких критериях сегментируются только самые светлые части аврорального свечения, а в ином случае почти все изображение.

Для свойства LPSS рассчитывается спектр области и окна при помощи прямого преобразования Фурье, сравнивая их можно сделать вывод о «размытии» окна. Так как изображение в окнах нормализовано относительно всего изображения, сравнение спектров окна и области дают верный результат. Для Gradient Histogram Span рассчитывается усредненный гауссиан, который является градиентом гистограммы окна и сравнивается с усредненным гауссианом области. Для Maximum Saturation используется просто значение интенсивности цвета в точке.

Для получения маски изображения на основе результата работы алгоритма используется функция Cv2.THRESH_OTSU(). Она преобразует изображение в бинарное (где есть только значение 255 или только значение 0), на основании гистограммы черно-белого изображения и минимальной среднеквадратичной ошибки вычисляется значение T и все пиксели исходного изображения, у которых значение больше T получают значение 255, а все остальные 0. Примеры результатов работы алгоритма видны на рисунках 13, 14, 15 (с маской).

Заключение

Полносвязная НС

В результате проведенной работы максимальная достигнутая точность нейронной сети приблизительно равна 0,6-0,7. Нейронные сети с наилучшей точностью использовали такие нелинейные функции активации как sigmoid и tanh, а количество итераций соответствовало 20-40. Т.е. полносвязная сеть достаточно хорошо показала себя в данной задаче классификации. Однако у нас есть такие типичные проблемы при обучении нейросетей как нехватка экспериментальных данных для обучения и затухающий градиент.

Из перспектив можно отметить увеличение количества экспериментальных данных, изменение модели нейросети или её типа, регуляризация весов.

Сверточная НС

Проанализированы методы и алгоритмы нахождения контуров, увеличение контрастности, построение структурных тензоров, создание масок и сегментации изображений. Реализован алгоритм для сегментации изображения полярного сияния и выделения на нем аврорального свечения с дальнейшим объединением, наложением и созданием маски. Реализована сверточная нейронная сеть, созданная на основе архитектур Inception и ResNet, а также используя блоки их реализующей модели InceptionResNetV2. Обучена модель разработанной архитектуры с максимальной точностью классификации 0,72 что для данного датасета является весомым и хорошим результатом. Максимальная точность получена с параметрами функций активации softmax и relu, learning_rate равный 0,0001, а batch_size равный 8.

В результате исследования показано, что системы на основе НС пригодны для распознавания таких сложных объектов как ПС. Несмотря на то, что полносвязные НС достаточно редко используются для распознавания сложных объектов на изображении, после проведения оптимизации в нашем случае полносвязная НС дала более высокую вероятность точной классификации по сравнению со сверточной моделью.

В качестве перспективы развития метода следует отметить, что архитектура сверточной сети позволяет исследовать изображения ПС в виде отдельных монохромных слоев формата RGB, HSB, HSL, CMYK, LAB и LCh, а также изображения зарегистрированные через светофильтры соответствующие свечению отдельных газов ПС.

Список использованных источников:

1. Лазутин Л.Л. Авроральная магнитосфера // Модель космоса / Под. ред. М.И. Панасюка, Л.С. Новикова М.: КДУ, 2007. Т. 1
2. Старков Г.В. Планетарная динамика аврорального свечения // Физика околоземного космического пространства. Апатиты: ПГИ, 2000. Т. 3.
3. Ф. Шолле – Глубокое обучение на Python. Питер, 2018 г.
4. Node.js Электронный ресурс]. – Режим доступа: <https://www.nodejs.org/>. – Дата доступа: 01.07.2022
5. ImageNet [Электронный ресурс]. – Режим доступа: <https://www.image.net.org/>. – Дата доступа: 01.07.2022.
6. Power spectra and distribution of contrasts of natural images from different habitats [Электронный ресурс]. – Режим доступа: <https://www.sciencedirect.com/science/article/pii/S0042698903004711>. – Дата доступа: 01.07.2022.

Поступила в редакцию 04.07.2022.

ИССЛЕДОВАНИЕ ВОЗМОЖНОСТИ ОЦЕНКИ КАЧЕСТВА МОТОРНОГО МАСЛА ПО ЭЛЕКТРОФИЗИЧЕСКИМ ПАРАМЕТРАМ ПРИ ЭКСПЛУАТАЦИИ АВТОМОБИЛЕЙ

С.С. Беленькая, В.С. Курило, Е.С. Максимович, БГУ

Рассмотрена возможность использования диэлектрических параметров для оценки качества моторного масла. Установлена взаимосвязь диэлектрических параметров моторного масла и его показателей качества при старении моторного масла для моторного масла 5W-30, 10W-40. Обзор методов оценки качества моторного масла.

Актуальность задачи. Моторные масла являются одним из основных функциональных элементов двигателя, определяющим эффективность и надежность его работы. Оно обеспечивает смазку всех движущихся деталей двигателя, покрывая их защитной пленкой, сокращая трение и износ, предохраняет детали двигателя от вредных отложений, коррозии, участвует в защите мотора от перегрева [1].

Совершенствование конструкции двигателей и повышение качества моторного масла позволяют обеспечивать заявленный ресурс и снижать интенсивность изнашивания узлов трения механизмов и систем. При эксплуатации моторных масел происходят физико-химические процессы, сопровождающиеся срабатыванием присадок, накоплением продуктов износа, что приводит к преждевременному исчерпанию потенциала физико-химических свойств масла и снижению ресурса силового агрегата.

По регламенту автопроизводителей менять масло в двигателе нужно каждые 10-15 тыс. километров, либо не реже раза в год. Параметры, приведенные в руководстве по эксплуатации каждого автомобиля, характерные для нового двигателя, держатся только определенное время, после чего расход масла начинает повышаться. Вот почему решение вопросов, связанных с оперативным контролем качества моторного масла, является достаточно актуальной задачей. Контроль автомобильных масел на соответствие требованиям нормативной технической документации является главным условием надежной и долговечной работы техники, так как результаты, получаемые при его проведении, с одной стороны, позволяют повысить эффективность применения масел, с другой – выявить некондиционные.

Основные подходы к диагностике масел. Измерение диэлектрических свойств моторного масла может дать информацию о параметрах, являющихся критическими при производстве и эксплуатации. На значение диэлектрической проницаемости влияет молекулярная структура вещества и его межмолекулярные взаимодействия. Определение связи между проницаемостью и перечисленными параметрами является одной из самых важных и интересных научных задач. Особый интерес представляет изучение молекулярных взаимодействий в моторных маслах, где они менее всего изучены. В работе [2] рассматриваются электрофизические параметры масла. Авторами были получены интересные

результаты для 4 различных образцов: 1. Все изученные моторные масла как дисперсные системы являются неполярными диэлектриками ($\epsilon' < 3$) с ионной проводимостью. Удельная электрическая проводимость для всех масел при $t = 40^\circ\text{C}$ находится в интервале 17–130 нСм/м, а к 100°C увеличивается до 0.18–1.12 мкСм/м. 2. Скорость уменьшения удельного электрического сопротивления всех образцов масел при повышении температуры больше по сравнению с темпом падения их вязкости, поэтому «правило Вальдена–Писаржевского» не выполняется.

Комплексные показатели качества масел (ПКМ) используются для более объективной и всесторонней оценки его состояния и построены по принципу комбинирования единичных или использования основных эксплуатационных характеристик двигателя. Для прогнозирования допустимой периодичности смены моторных масел в двигателе используются браковочные нормы как информативные показатели качественного состояния (например процесса образования высокотемпературных отложений [3]). В частности, предлагается осуществлять оценку состояния моторного масла по ОКП – обобщенному комплексному показателю, который состоит из суммы единичных диагностических показателей [4].

Из рассмотренных в [5] методов ПКМ видно, что комплексные ПКМ формируемые с использованием физических принципов (определение оптической плотности масла) обладают такими же недостатками, как и комплексные ПКМ построенные по принципу комбинирования единичных показателей. Они также не пригодны для оперативного определения качества масел в полевых условиях, условиях мастерских автотранспортных предприятий и станций технического обслуживания, по причине высокой трудоемкости определения каждого показателя входящего в их расчетные уравнения и необходимости применения дорогостоящего стационарного оборудования для их определения. Так, например, существенным недостатком методики, приведенной в работе [6] является тот факт, что ее применение базируется на некоторых допущениях:

- во-первых, принимается, что «угоревшее» масло непрерывно компенсируется доливкой свежего с той же скоростью, за счет чего количество масла в системе смазки сохраняется постоянным;
- во-вторых, принимается, что загрязнения равномерно распределены во всем объеме работающего мас-

ла и расходуются с «угорающим» маслом, имеющим ту же концентрацию загрязнений.

Как показывает практика, у рассмотренных комплексных показателей качества отсутствуют реально обоснованные браковочные значения, при достижении которых эксплуатацию масла в двигателе необходимо прекратить. Кроме того, все комплексные ПКМ полностью не пригодны для оперативного определения качества масел в полевых условиях, условиях мастерских автотранспортных предприятий и станций технического обслуживания. Это связано с трудоемкостью определения каждого единичного показателя, и необходимости применения дорогостоящего стационарного оборудования для их определения.

Некоторые комплексные ПКМ формируются с использованием физических принципов, отличных от применяемых при оценке единичных показателей (в частности определение оптической плотности масла). Например, авторы работы [7] предлагают осуществлять комплексную оценку качества масла с использованием оптической плотности на лабораторной установке высокотемпературного каталитического окисления. Испытуемое масло (100 г) заливают в стальные стаканы и при температуре 2300 С, в течении 3 ч перемешивают медными стержнями, вращающимися с частотой 6000 мин⁻¹. В условиях испытания происходят процессы термодинамической деструкции и в меньшей степени окисления. Образующаяся в результате этих процессов дисперсная фаза коагулирует, что приводит к росту оптической плотности и вязкости масла. Рост этих параметров свидетельствует о накоплении предшественников нагарообразования (оптическая плотность) и степени окисления (вязкость). Такой контроль является достаточно трудоемким и долгим.

При эксплуатации моторного масла происходит возрастание показателя диэлектрической проницаемости. Это часто связано с окислением углеводородов масла, накоплением продуктов износа, добавлением присадок, других загрязнений, которые имеют высокую удельную электропроводность. Показатель диэлектрической проницаемости для разных типов масла может значительно отличаться. Разнородный состав смазочных материалов значительно затрудняет разработку единой методики определения качества масла относительно одного критерия. Критерий, который условно был принят за универсальный, получил свое определение в выражении, качественно описывающем влияние проводящих примесей на величину параметра $\Delta\epsilon$ моторного масла:

$$\Delta\epsilon = \left[\left(\frac{1}{1 - \frac{1}{3} \frac{\rho_{\text{м.}}}{\rho_{\text{п.п.}}} \omega m} \right) - 1 \right] \epsilon, \text{ где}$$

$\rho_{\text{м}}$ – плотность моторного масла г/см³;

$\rho_{\text{п.п.}}$ – плотность проводящих примесей находящихся в масле г/см³;

ωm – относительная массовая концентрация проводящей примеси,

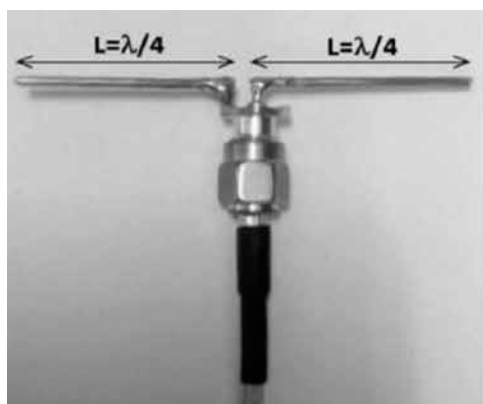
ϵ – диэлектрическая проницаемость полярных диэлектриков при небольших электрических полях и температурах не ниже комнатной.

Значение показателя $\Delta\epsilon$ может служить фактором определения качества масла, но индивидуально для каждого двигателя, учитывая его техническое состояние, качество применяемого масла, условия эксплуатации, сезон использования, а также временной ресурс хранения. Если масло окисляется [8, 9] и двигатель неправильно смазан, его может заклинить. Известно, что масло окисляется, даже если оно не используется определенное время [10]). Сложность заключается в том, что нет универсального показателя, указывающего, когда необходимо заменить машинное масло. Повышения информативности методов определения качества моторного масла по его диэлектрическим параметрам можно добиться разработкой методики измерения значения диэлектрической проницаемости на разных частотах, с установлением степени влияния отдельных показателей моторного масла на значение ϵ на каждой из измеренных частот [11]. Существующие промышленные встроенные датчики масла служат только для измерения уровня в баке и его температуры, но никогда не информируют водителя о качестве моторного масла. Если замена масла не произведена вовремя, последствия могут быть очень серьезными для транспортного средства, что может привести к повреждению двигателя [12].

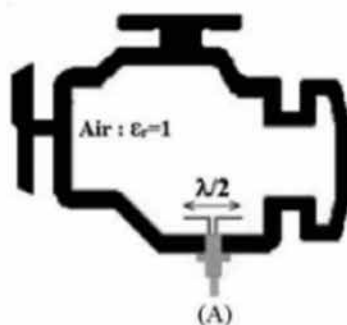
Разработка датчиков, позволяющих производить оценку состояния масла, является достаточно важной задачей. Датчики такого плана можно комбинировать с системой RFID [13], передающей беспроводным путем информацию. Это позволит получить сигнал о точном моменте, когда надо масло менять. Эта идея может внести значительный вклад в оптимизацию расходов на техническое обслуживание, особенно в государственных автомобильных парках. Оптимизация срока службы моторного масла, а также масляного фильтра сэкономит не только средства, но и время, трудозатраты, увеличит срок службы, эффективность и защитит двигателя.

Одним из вариантов таких систем, предлагаемых в работе [14] может быть датчик, состоящий из полувольтового диполя, помещенного внутрь масляного поддона автомобиля (рисунок 1). Оценивая изменение относительной диэлектрической проницаемости моторного масла, в зависимости от времени работы двигателя, можно оценить время необходимой замены. На рисунке 1 показан внешний вид датчика, реализующего данный метод и вариант его размещения.

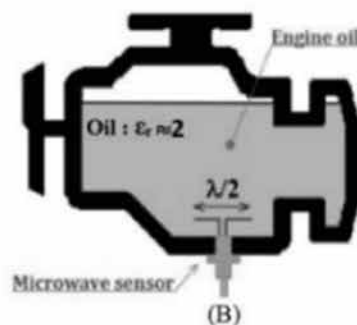
В данной работе резонансная частота диполя в свободном пространстве составляла 2,4 ГГц. В последствии, диполь, состоящий из плечей длиной по $L=29,8678$ мм ($\approx \lambda_0/4$) (рисунок 2), помещается в однородную диэлектрическую среду. Первичный съем результатов производился в пространстве без масла, затем – с масляным заполнением. Возбуждение произ-



а)



(A)



(B)

б)

Рисунок 1 – Система оценки изменения диэлектрической проницаемости машинного масла.

а) – внешний вид датчика, б) вариант размещения датчика

водилось через стандартный 50-Омный коаксиал с *sma*-разъемом. Результаты измерений были выполнены с использованием векторного анализатора цепей Agilent Technologies E5062A. Авторами показано тестирование трех проб – чистого масла, и после пробега 5 000 и 10 000 км. В качестве объекта исследования было выбрано масло QUARTZ 9000 – 5W40. Резонансный пик диполя в свободном пространстве был на частоте 2.4 ГГц, при погружении в кювету с маслом – на частоте 1.94 ГГц. Зная эти данные, можно оценить диэлектрическую проницаемость масла. Далее авторы приводят зависимости фазы и мнимой части входного импеданса детектора от частоты, полученные с использованием векторного анализатора и демонстрирующие чувствительность резонансной частоты детектора изменениям в масле в зависимости от пробега автомобиля (Она варьируется в зависимости от степени деградации масла). Смещение резонансного пика по частоте выглядело следующим образом. Для пустой емкости резонанс был на частоте 2.4 ГГц ($\epsilon=1.09$), для емкости, заполненной чистым маслом – на частоте 1.94 ГГц ($\epsilon=1.667$), для емкости, заполненной маслом после пробега 5000 км – на частоте 1.91 ГГц ($\epsilon=1.719$) и после пробега 10000 км – на частоте 1.89 ГГц ($\epsilon=1.765$).

Эксперимент и обсуждение. Данная идея была апробирована в лабораторных условиях на кафедре радиофизики и цифровых медиа технологий факультета радиофизики и компьютерных технологий БГУ. В качестве сенсорной системы была применена конструкция полуволнового резонатора, внешний вид которого показан на рис.2. К стандартному разьему N – типа с волновым сопротивлением 50 Ом была присоединена медная трубка с внутренним диаметром 5 мм и центральным проводником 1,3 мм. В трубке сделан продольный разрез длиной 36 мм. В начале и в конце разреза центральный проводник соединен медной проволокой диаметром 0,5 мм с двусторонней пайкой. В качестве излучателя и приемника использовался векторный анализатор цепей Agilent 8722 ET (50MHz – 40 MHz). В качестве объекта исследований взяты три образца масла:

5W-30 (образец 1 – вязкость климатическая при 100С равна 15,28 мм²/с, индекс вязкости 157, массовая доля механических примесей 0,015%, 10W-40 (образец 2 – вязкость климатическая при 100°С равна 15,28 мм²/с, индекс вязкости 157, массовая доля механических примесей 0,015%, 10W-40 (образец 2 – вязкость климатическая при 100°С равна 12,26 мм²/с, индекс вязкости 172, массовая доля механических примесей отсутствует), и одно отработанное масло (образец 3).



Рисунок 2 – Внешний вид коаксиального датчика для контроля диэлектрических свойств машинного масла.

На рисунке 3 приведены результаты исследований: резонансная частота датчика составила 2,6 ГГц при измерении в свободном пространстве, при погружении датчика в масло (образец 1) частота сдвинулась до 1,894 ГГц, для образца 2 резонансная частоты составила 1,873 ГГц, а при погружении в отработанное масло (образец 3) частота достигла 1,844 ГГц.

Заключение

Результаты исследований показывают, что СВЧ методы чувствительны даже к незначительным изменениям свойств жидких материалов и могут быть хорошей основой для разработки методов экспресс диагностики в промышленности и сельском хозяйстве.

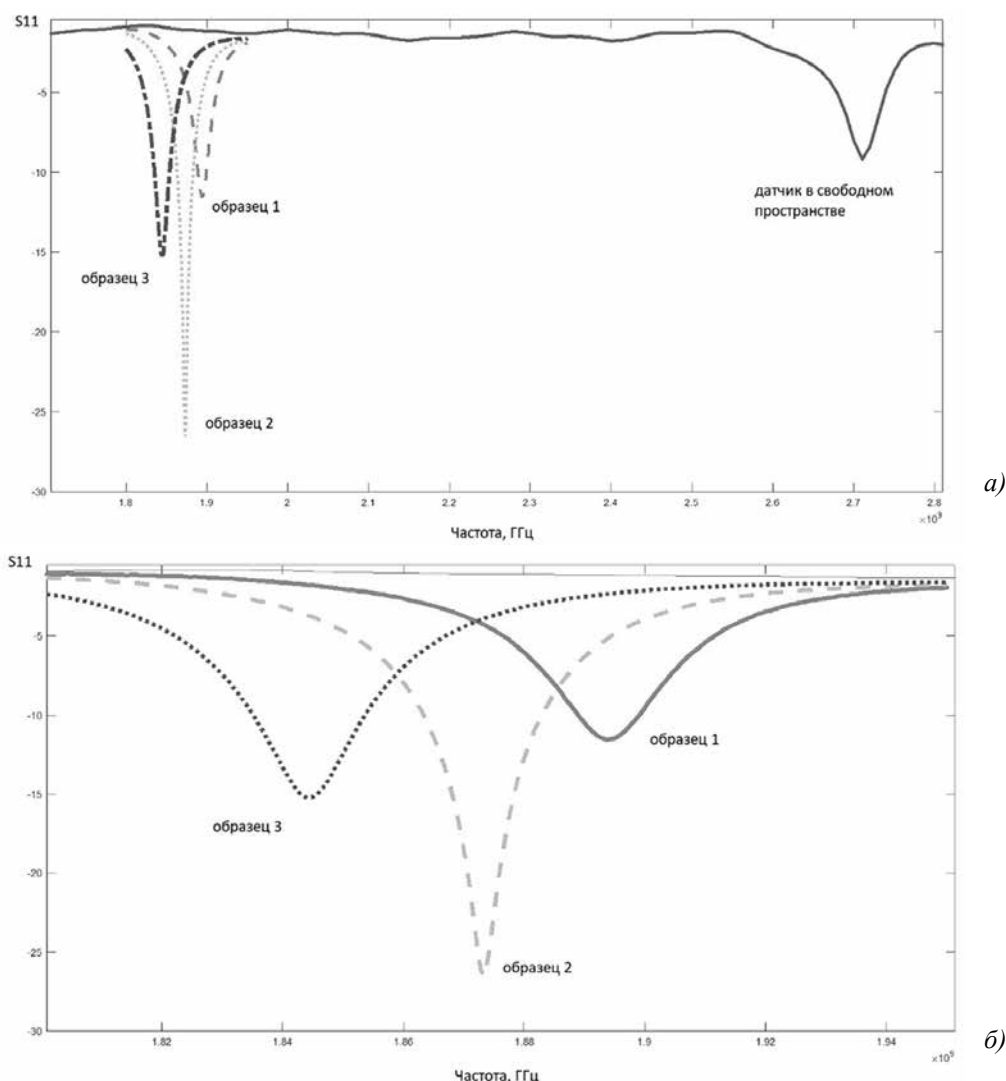


Рисунок 3 – Результаты измерений для трех разных типов масла. а) – разница между характеристиками датчика в свободном пространстве и при погружении в масло; б) – значения резонансных частот для трех образцов масла.

Литература

1. Alessandra Borin, Ronei J. Poppi. Application of mid infrared spectroscopy and iPLS for the quantification of contaminants in lubricating oil. Elsevier. Vibrational Spectroscopy 37(2005) 27-32.
2. И. В. Измestьев, О. А. Кожевникова, Поляризация и электропроводность некоторых моторных масел как дисперсных систем // Вестник Пермского государственного национального исследовательского университета Серия Физика, 2016, т. 32, стр. 43-51.
3. Золотов В.А. Научно-методические основы прогнозирования периодичности смены моторных масел в двигателях // Трение, износ, смазка. – Т. №10, – 2008. – № 3. – С. 71-79.
4. Лышко Г.П. Оптимизация сроков замены моторных масел // Химия и технология топлив и масел. – 1982 – № 11 С. 32–34.
5. А.Б. Григоров, И.С. Наглюк, «Рациональное использование масел», Харьков, Издательство «Точка», 2013, 183 стр.
6. В.Д. Резников, Э.Г. Синайский, Расчет концентрации не растворимых загрязнений в работающем масле // Двигательное строительство. – 1983. – №8. – С. 41–42.
7. А.Я. Левин, Г.Л. Трофимова, О.В. Иванова, Г.А. Будановская, Новые лабораторные методы оценки качества моторных масел // Химия и технология топлив и масел. – 2006. – № 2. – С. 50–5122.
8. P. Vahaoja. Oil analysis in machine diagnostics. ACTA Universitatis Ouluensis. A Scientiae Rerum Naturalium 458 (2006).
9. F. Bowman, G.W. Stachowiak. Determining the oxidation stability of lubricating oils using sealed capsule differential scanning calorimetry. Elsevier. Tribology International 29 (1996) 27-34.
10. Benzing R, Goldblatt I, Hopkins V, Jamison W, Mecklenburg K, Peterson M. Friction and wear devices. American Society of Lubrication Engineers (1976).
11. А.А. Хазиев, Н.Н. Сугатов, М.Ю. Петухов, Оценка качества моторного масла при эксплуатации автомобилей диэлектрическими методами // Модернизация и научные исследования в транспортном комплексе, 2014 г., стр. 223-228.
12. Gautam, M., Chitoor, K., Durbha, M., & Summers, J. C. (1999). Effect of diesel soot contaminated oil on engine wear — investigation of novel oil formulations. Tribology International, 32(12), 687–699.
13. Bakker, A., & Huijsing, J. H. (1996). Micropower CMOS temperature sensor with digital output. IEEE Journal of Solid-State Circuits, 31(7), 933–937.
14. Mejri, F., & Aguilu, T. (2017). Design of a passive microwave sensor for the characterization of mobile engine oil. 2017 IEEE-APS Topical Conference on Antennas and Propagation in Wireless Communications (APWC).

Поступила в редакцию 04.07.2022.

ЛАЗЕРНОЕ ВОЗДЕЙСТВИЕ НА ТКАНЬ 07С11-КВ С КОМБИНИРОВАННЫМ ПОКРЫТИЕМ ЦИРКОНИЯ И УГЛЕРОДА

УДК 539.2

М.И. Маркевич¹, В.И. Журавлева², Т.Ж. Кодиров³,
И.П. Акула¹, Н.М. Чекан¹, Е.Н. Щербакова¹
¹ГНУ «Физико-технический институт НАН Беларуси»
²Военная академия Республики Беларусь
³Ташкентский институт текстильной и легкой промышленности

Взаимодействие лазерного излучения с композиционными материалами на основе тканых материалов является основой новых технологических процессов.

В работе исследована морфология поверхности композиционного материала на основе смесовой ткани 07С11-КВ (Моготекс) с покрытием углерода и циркония после лазерного воздействия в режиме сдвоенных импульсов. Установлены режимы лазерного воздействия, при которых происходит разрушение композиционного материала в результате расплавления материала под действием концентрированного потока лазерного излучения. Установлено, что при увеличении вложенной энергии до 15 Дж происходит перфорация ткани в результате расплавления материала под действием концентрированного потока лазерного излучения.

Введение

В настоящее время важной задачей является разработка комбинированных текстильных материалов. В последнее время в государствах СНГ расширяется ассортимент технического текстиля для различных отраслей промышленности, для обувных изделий и специальной одежды [1]. Актуальность в настоящее время приобретает защитный текстиль. Тканые экранирующие материалы широко применяются в различных отраслях промышленности [4-6]. Металлизацию тканей часто применяют на синтетических полиэфирных полотнах. Известно, что повышение электрической проводимости ткани можно достичь при создании тканей путем введения металлизированных нитей при производстве тканей, нанесением покрытия толщиной до 15 мкм методом вакуумного нанесения [1].

Для изготовления специальной защитной и медицинской одежды используются смесовые ткани, содержащие полиэфирные и хлопковые волокна в определенных соотношениях. Основным недостатком смесовых тканей является повышенная склонность к электризации. Так, например, возникновению статических зарядов на одежде медицинского персонала способствует значительное количество электроприборов в современных диагностических кабинетах, низкая влажность воздуха, фоновый уровень ионизирующего излучения, что вызывает у человека нарушение обмена веществ, работы сердечнососудистой системы, утомляемости [2,3].

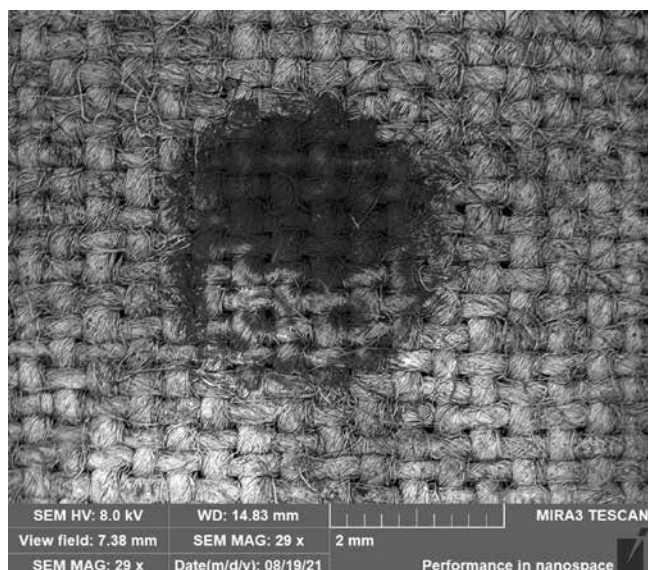
Данные как зарубежных, так и отечественных исследователей свидетельствуют о влиянии электромагнитных полей (ЭМП) во всех частотных диапазонах на живые организмы. Согласованность деятельности органов и систем обусловлена колебательными процессами в клетках и различных органах, что усиливает биологическую активность от ЭМП. [2, 3]. Дополнительно к изложенному, особую значимость приобретают ткани для одежды тепловой маскировки, что составляет наиболее сложную задачу, так как надо соблюсти термодинамическое равновесие биологического объекта с окружающей средой. С целью усиления теплоотвода для эффективности тепловой маскировки эффективно использовать лазерную перфорацию тканей. Поэтому исследования в этом направлении являются актуальными.

Целью работы является исследование изменения морфологии поверхности ткани 07С11-КВ с комбинированным покрытием циркония и углерода при лазерном воздействии в режиме сдвоенных импульсов.

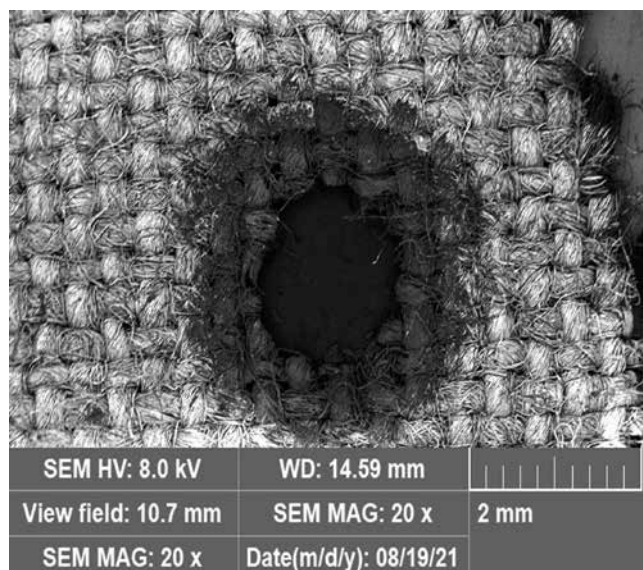
Основная часть

В данной работе пленки циркония осаждались в вакууме с использованием источника стационарной металлической плазмы (цирконий), работающего в режиме сепарации. Покрытие циркония наносилось на смесовую ткань 07С11-КВ (производство «Моготекс»). Предварительно перед формированием покрытий поверхность тканей обрабатывалась высокоэнергетичными ионами аргона для удаления органических загрязнений в течение 15 минут при следующих параметрах: давление аргона в вакуумной камере порядка $3,7 \times 10^{-2}$ Па, ускоряющее напряжение 2000 В, ионный ток 20 мА. Затем включался источник катодно-дуговой плазмы, установленный в режиме электромагнитной сепарации потока на 900 и выполнялось осаждение покрытия металла катода. Поскольку температура покрытия при его формировании на поверхности основы может достигать нескольких сотен градусов Цельсия, то процесс велся путем чередования периодов работы источника плазмы (1 минута) и паузы для охлаждения ткани (1 минута). Для осаждения углерода использовались два источника дуговой плазмы (напряжение разряда 300 В, частота следования импульсов 3 Гц, число разрядных импульсов 15000).

Исследование структуры синтезированных материалов производилось с использованием растрового электронного микроскопа MIRA 3.



а)



б)

Рисунок 1 – Структура ткани после лазерного воздействия: а) вложенная энергия 2 Дж, время воздействия 2 с, б) вложенная энергия 15 Дж, время воздействия 15 с.

Исследования элементного состава образцов проводились с помощью системы энергодисперсионного (EDS) микроанализа, установленной на сканирующем электронном микроскопе MIRA3.

Для обработки материала использован лазер на алюмоиттриевом гранате (LS-2134D) с длиной волны 1064 нм, генерирующий в двухимпульсном режиме (импульсы разделены временным интервалом 3 мкс, длительность импульсов 10 нс, частота следования импульсов 10 Гц, энергия одиночного импульса ~0,05 Дж).

Образованная в результате испарения вещества под действием первого импульса абляционная плазма создает в приповерхностном слое область с повышенной температурой и пониженной плотностью частиц, что приводит к более полному использованию энергии второго импульса для лазерной абляции [3-5]. Образец облучали лазерным излучением в интервале энергий 2-15 Дж при временах экспозиции от 2 до 15 секунд. Размеры образца: толщина ~0,5 мм, длина – 15 мм, ширина – 20 мм.

Пороги плазмообразования большей частью зависят от качества покрытия, неоднородности материала, структуры покрытия, дефектов, состава покрытия. Характер световой эрозии композиционного материала зависит от особенностей самого материала: его оптических (соотношение коэффициентов пропускания, поглощения и отражения света), теплофизических, структуры и других характеристик [5].

Механизмы лазерного разрушения тканых материалов также зависят от их строения и структуры покрытий и сильно отличаются друг от друга [7,8].

Следует отметить, что данный класс композиционных тканых материалов не является изученным, поэтому экспериментальные результаты в этом на-

правлении позволят определять параметры порога разрушения этих материалов. На рисунке 1 представлена морфология данного материала после лазерного воздействия.

Из рисунка 1а) следует, что при данных условиях лазерного воздействия не происходит перфорации ткани. На рисунке 1а) показано, что взаимодействие излучения с поверхностью покрытия и его удаление



Рисунок 2 – Структура ткани после лазерного воздействия (на краю зоны воздействия): а) вложенная энергия 2 Дж, время воздействия 2 с

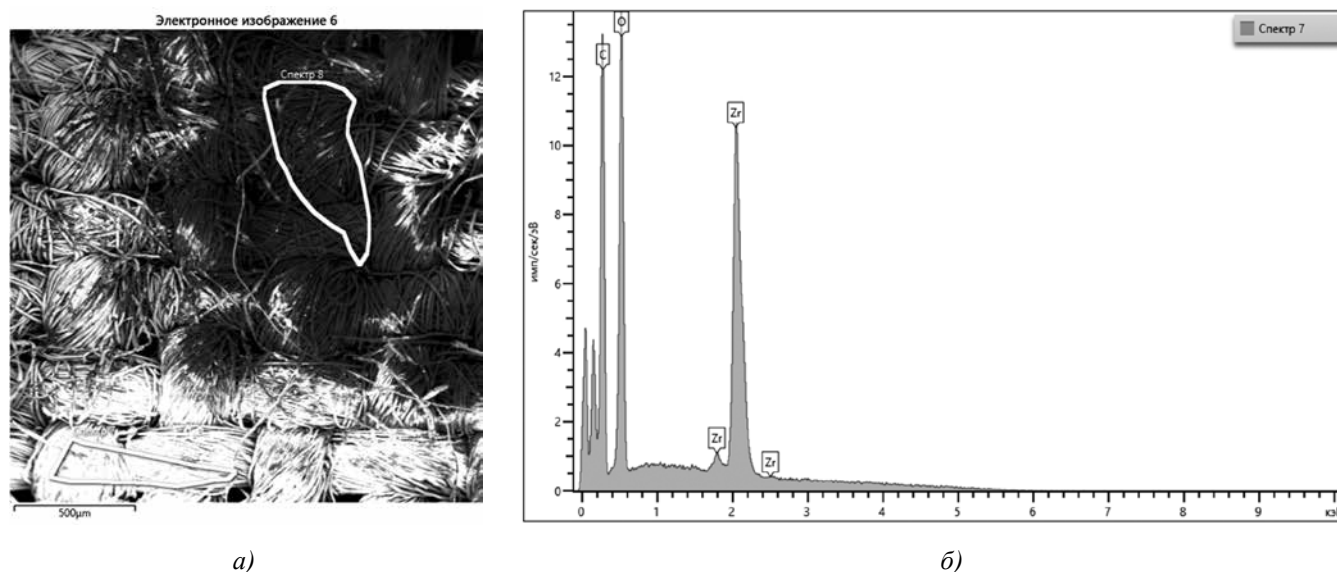


Рисунок 3 – Морфология поверхности и элементный анализ ткани с покрытием после лазерного воздействия (вложенная энергия 15 Дж, время воздействия 15 с): а) морфология поверхности; б) элементный анализ

произошло не по всей поверхности пятна воздействия. Диаметр пятна составляет ~3 мм. Вынос покрытия произошел с пятна диаметром 1,7 мм. При увеличении вложенной энергии до 15 Дж зона воздействия расширяется и происходит перфорация ткани.

Диаметр отверстия составляет 1,6 мм. Вокруг отверстия наблюдается зона выноса покрытия. Размер зоны составляет 0,9 мм. При данных вложенных энергиях (10-15 Дж) происходит плавление ткани с комбинированным покрытием, и образуются сквозные отверстия. Форма нитей вблизи очага плавления вследствие высокой температуры искажается, и они приобретают бугорчатый вид. При таких режимах лазерного воздействия реализуются условия образования низкотемпературной плазмы, температура при этом значительно превышает температуру термической деструкции.

На рисунке 2 представлена морфология поверхности ткани на краю зоны воздействия.

Как следует из рисунка 2 покрытие на поверхности ткани осталось, но однородность покрытия сильно нарушена, оно имеет пятнистый характер. На рисунке 3 представлена морфология поверхности и элементный анализ ткани с комбинированным покрытием циркония и углерода при вложенной энергии 15 Дж и длительности 15 с.

Элементный анализ, проведенный по площади областей, фиксирует состав нанесенного покрытия (цирконий и углерод).

Выводы

Диагностирована морфология поверхности смесовой ткани 07C11-KB (Моготекс) с комбинированным покрытием циркония и углерода после лазерного воздействия в двухимпульсном режиме при вложенной

энергии 2-15 Дж и времени экспозиции 2- 15 секунд. Показано, что при вложенных энергиях до 7 Дж не происходит перфорации ткани, а происходит удаление покрытия без нарушения структуры тканого материала. Установлено, что при увеличении вложенной энергии до 15 Дж и времени воздействия до 15 с происходит перфорация ткани в результате расплавления материала под действием концентрированного потока лазерного излучения, рядом с отверстием наблюдается зона выноса покрытия до 800 мкм. Созданы физические основы для технологии лазерной перфорации композиционных материалов на основе циркония и углерода на смесовых тканях, что является базой новых технологических процессов.

Литература

1. Бондарчук, М. М. Подходы к классификации технического текстиля / М. М. Бондарчук // Проблемы современной науки и образования. – 2015. – № 11. – с. 95–99.
2. Пряхин, Е.А. Влияние неионизирующих электромагнитных излучений на животных и человека/ Е.А. Пряхин, А.В. Аксеев//Челябинск.- Полиграф-Мастер.– 2006. – С.220.
3. Тихонов, М.Н. Прогрессирующая трансформация электромагнитной среды обитания человека: медико-экологическая проблема XXI века/ М.Н. Тихонов, В.В. Довгуша // Экол. системы и приборы. –2000. – №1. – с.46–54.
4. Вейко, В.П., Метев, С.М. Лазерные технологии в микроэлектронике /В.П. Вейко, С.М. Метев // София.- Болгарская Академия Наук. – 1991. – С.358.
5. Маркевич, М.И. Структурные превращения в тонких металлических пленках при импульсном лазерном воздействии/М.И. Маркевич, А.М. Чапланов

// Известия Национальной академии наук Беларуси. – 2016. – №1. – с.28–34.

6. Ласковнев, А.П. Воздействие лазерного излучения на композиционные материалы на основе лавсановых тканей и углерода/ А.П. Ласковнев, А.Г. Анисович, Маркевич, В.И. Журавлева, Н.М. Чекан // Материалы конф. «Композиционные материалы на основе техногенных отходов и местного сырья: состав, свойства и применение». – 2021. –Ташкент. – с. 30–31.

7. Анисович, А.Г. Воздействие лазерного излучения на лавсановую ткань, покрытую углеродом / А.Г. Анисович, В.Г. Залесский, М.И. Маркевич, А.Н. Малышко, В.И. Журавлева, Н.М. Чекан, Н.М., Чэнь Чао // Полимерные материалы и технологии. – 2020. – Т.6. – №1. – с. 72–77.

8. Воробьев, В.С. Плазма, возникающая при взаимодействии лазерного излучения с твердыми мишенями / В.С. Воробьев// УФН. 1993. Т. 163(12) . с. 51–83.

Abstract

The interaction of laser radiation with composite materials based on woven materials is the basis of new technological processes.

The morphology of the surface of a composite material based on a 07C11-KV (Mogotex) mixed fabric coated with carbon and zirconium after laser exposure in the mode of double pulses is investigated. The modes of laser exposure are established in which the destruction of the composite material occurs as a result of the melting of the material under the action of a concentrated laser radiation flux. It is established that when the energy invested increases to 15 J, tissue perforation occurs as a result of melting of the material under the action of a concentrated laser radiation flux.

Поступила в редакцию 31.05.2022 г.

МНОГОПОЛЬЗОВАТЕЛЬСКАЯ СИСТЕМА УДАЛЕННОЙ 3D FDM ПЕЧАТИ

УДК 004.77

Б.С. Петкевич, И.А. Романов
Белорусский государственный университет

Введение

Аддитивные технологии (АТ) являются одним из наиболее стремительно развивающихся направлений за последнее десятилетие, наряду с искусственным интеллектом, квантовыми вычислениями, автономными роботами [1]. В настоящее время трехмерная печать применяется в автомобильной промышленности для создания прототипов и мелкосерийного производства деталей [2], медицине для печати биосовместимых имплантатов и костных тканей [3], аэрокосмической отрасли для печати двигателей ракет, их составных частей [4], а так же для печати элементов турбореактивных двигателей, элементов летательных аппаратов и искусственных спутников земли [5]. В зависимости от требований, предъявляемых к физическим свойствам деталей, применяются различные материалы и технологии трехмерной печати: моделирование методом наплавления расплавленного материала [6], селективное лазерное плавление и спекание порошков пластмасс и металлов [2], струйная печать методом послойного нанесения связующего материала [7], фотополимеризация. Преимуществами аддитивных технологий перед традиционными методами литья и механической обработки являются экономия материала, создание сложных форм и полостей деталей не достижимых другими методами производства, а так же возможность создания объектов из нестандартных для

производства материалов, таких как биологические ткани [3], высокотемпературные градиентные сплавы CuNi [8], угленаполненные композиты [9]. Основным недостатком АТ является низкая скорость производства деталей [10]. Кроме того, детали, получаемые методами селективного лазерного плавления, фотополимеризации или наплавления имеют шероховатую поверхность, могут подвергаться возникновению внутренних напряжений и усадке в процессе печати.

Самой распространенной технологией 3D печати является моделирование методом наплавления (fused deposition modeling – FDM) [6]. Технология FDM разработана Скоттом Крампом в 1988 г. Благодаря своей низкой стоимости эта технология приобрела широкую популярность среди любителей 3D печати. В качестве материалов FDM печати, как правило, применяются прутки термопластичных полимеров (ABS, PLA, PC, PET, PMMA, воск и др.), а так же композиты – термопластики, наполненные углеродными нанотрубками, металлическими частицами, угле- и стекловолокном [6, 9].

В конце 2021 года компания «3D Printing Industry» провела опрос среди ученых, руководителей высшего звена и инженеров крупных компаний о будущем развитии аддитивных технологий в ближайшие годы. В связи с возникшими проблемами в цепочках поставок комплектующих и задержкой в логистике прогнозируется массовое внедрение технологий 3D печати в

производственные процессы в ближайшие годы [11]. Представители компаний «Materialise» и «Nexa3D» отмечают тенденцию к использованию автоматизации и контроля качества процессов аддитивного производства с применением технологий машинного обучения и машинного зрения [11].

Основными трудностями массового внедрения АТ в производство являются невысокая скорость 3D печати, необходимость в наличии технологии и квалифицированных специалистов, высокие требования к качеству деталей и их сертификация [12]. Для производства готовой детали по технологии FDM оператору необходимо выполнить следующие действия: преобразовать исходную модель в исполнительный код (g-файл), отправить g-файл на печать, по окончании печати извлечь деталь и произвести контроль качества. Удаленное управление и автоматизация процессов FDM печати позволит существенно сократить время и стоимость производства продукции и создания прототипов, а так же позволит избавиться от оператора и исключить человеческий фактор [12].

На сегодняшний день существуют компании, которые добились успеха в автоматизации процесса обработки 3D моделей без использования дополнительного программного обеспечения. Оборудование компаний «Hewlett-Packard» и «Desktop Metal» имеют встроенный слайсер для нарезки модели на слои. Однако по-прежнему не существует полноценной технологии автоматизированного аддитивного производства, которая смогла бы обходиться без участия человека. Нидерландский стартап Blackbelt3D представил первую модель конвейерного FDM принтера в 2017 году [13]. Областью построения модели является конвейерная лента на основе карбона. Экструдер перемещается в плоскости, расположенной под углом 45° к области построения модели. В 2019 году компания «Creality» представила конвейерный принтер «Creality CR-30» с такой же кинематической схемой. Недостатками этих принтеров являются наличие открытой области печати, недостаточная адгезия материалов к конвейерной ленте [14], нестандартная кинематическая схема, требующая применения слайсера собственной разработки, отсутствие аппаратной и программной частей, позволяющих реализовать удаленное управление процессом 3D печати.

Таким образом, актуальной задачей является разработка автоматизированных систем аддитивного производства с функцией удаленного контроля и управления, позволяющих ускорить и удешевить внедрение АТ в промышленность.

Целью настоящей работы является разработка многопользовательской системы удаленной 3D печати на базе медиасервера и FDM принтера с конвейерной лентой. Основные требования, предъявляемые к системе: возможность непрерывной печати, возможность постановки модели в очередь на печать удаленно, удаленное управление параметрами печати и наблюдение за процессом печати.

Реализация аппаратной части

Для обеспечения непрерывного процесса печати деталей разработан FDM 3D принтер с конвейерной лентой, представленный на рисунке 1. Принтер имеет область построения модели и зону охлаждения детали. Построение модели осуществляется на конвейерной ленте. Экструдер перемещается в плоскости, параллельной области построения, называемой порталом. Портал имеет стандартную кинематическую схему «Core XY». Перемещение портала осуществляется по вертикальной оси ‘Z’ благодаря чему становится возможным применение стандартных слайсеров для нарезки модели на слои. Когда печать первой детали закончена, конвейерная лента перемещает готовую деталь в зону охлаждения. В процессе охлаждения первой детали принтер может начать печать следующей. Таким образом, обеспечивается непрерывный процесс печати без участия оператора. Благодаря тому, что после окончания печати текущей детали печать следующей начинается через ~30 с, а система может работать непрерывно 24 часа в сутки удастся существенно увеличить производительность 3D печати по сравнению с традиционным типом 3D печати, в котором принимает участие оператор.

Управление процессом печати осуществляется посредством одноплатного компьютера Raspberry Pi с прошивкой Klipper на базе операционной системы Linux. Управление периферией (шаговые двигатели, нагревательные элементы, датчики температуры) осуществляется платой MKS GEN. Klipper имеет свой веб-интерфейс, позволяющий взаимодействовать пользователю с принтером внутри сети. Веб-интерфейс содержит настройки параметров печати, отображение температур экструдера и зоны построения, возможность управления параметрами печати в процессе (изменение температуры, скорости печати), возможность загрузки g-файлов и запуска печати. Так как Klipper настроен на работу с одним пользователем и в нем отсутствует функция автоматической печати с постановкой в очередь необходимо применение дополнительных модулей, позволяющих реализовать требуемый функционал. Наиболее удобным способом взаимодействия пользователя с системой удаленной 3D печати является веб-приложение.

Аппаратная часть многопользовательской системы удаленного управления процессом 3D печати состоит из связанных между собой физических модулей:

Конвейерный 3D принтер с управляющей платой Raspberry Pi;

Медиасервер для взаимодействия пользователя с интерфейсом веб-приложения и управления процессом печати. Этот блок реализован на отдельном компьютере Raspberry Pi;

IP-видеокамеры.

В разработанной системе реализована возможность подключения нескольких конвейерных принтеров, ос-

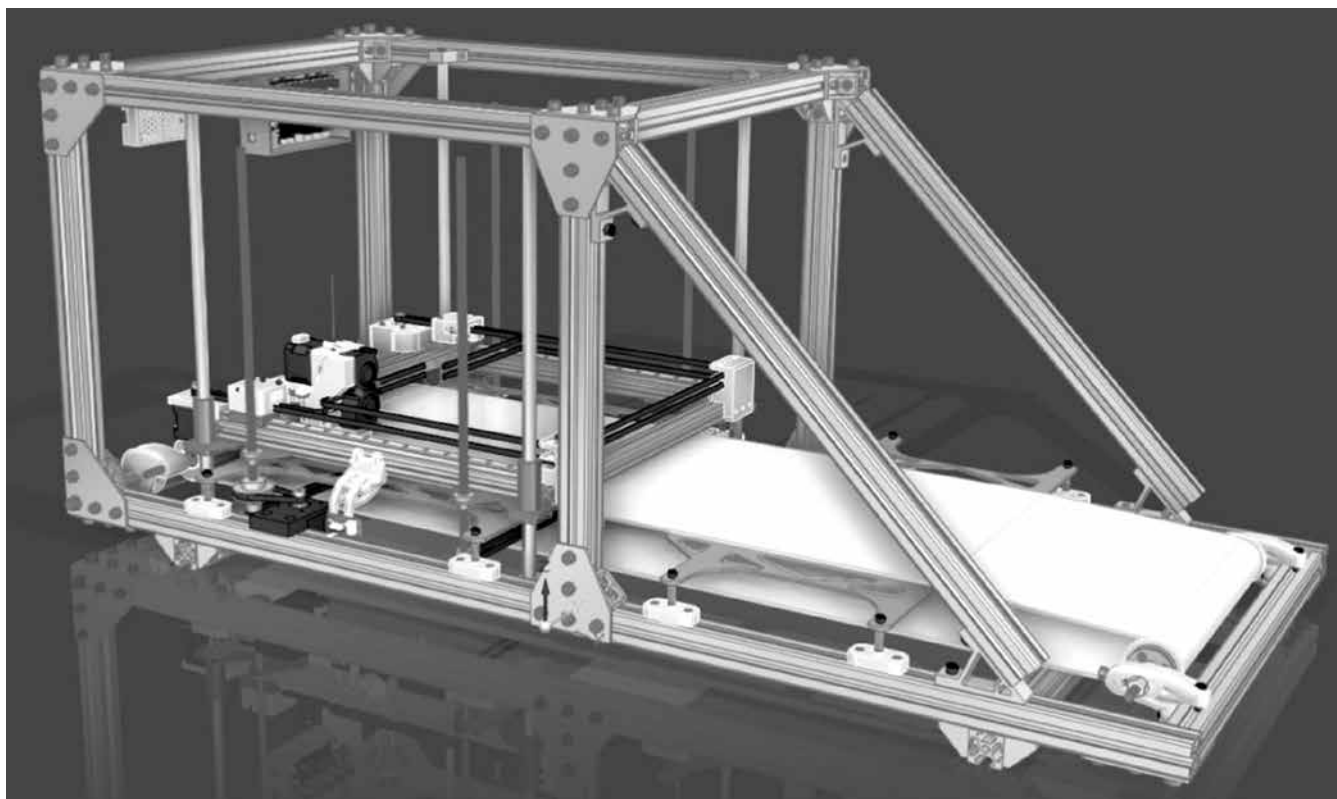


Рис. 1. Разработанный конвейерный принтер

нащенных платой Raspberry Pi с прошивкой Klipper. Аппаратное разделение процесса печати с процессами взаимодействия пользователей с системой, хранения и обработки данных позволяет повысить отказоустойчивость и производительность системы.

Панель управления системой

Управление принтером осуществляется через окно веб-приложения. Для получения доступа к принтеру необходимо пройти авторизацию как пользователь системы либо как администратор. После авторизации пользователю доступна панель управления с вкладками: “добавить новый заказ”, “видеонаблюдение”, “история заказов”, изображенная на рисунке 2. Пользователю доступны следующие функции:

- Загрузка файла модели на печать (g-код), выбор материала и цвета;

- Наблюдение за процессом печати только своей модели, получение информации о статусе печати;

- Остановка печати;

- Получение информации о начале и окончании печати путем оповещения на e-mail. По окончании печати на e-mail пользователя отправляется изображения готовой модели;

- Получение оценочного времени о начале и окончании печати, получение информации об ошибках;

- Функция вызова техподдержки.

После загрузки файла модели система передает задачу на печать в очередь. Так как приблизительное

время печати каждой модели в очереди известно, итоговое время начала печати следующей модели определяется суммой времён печати каждой предыдущей модели, стоящей в очереди на печать. Если очередь пуста, печать начинается сразу после отправки пользователем файла и обработки его системой.

Для администратора кроме вкладок пользователя доступна вкладка “статус принтеров”. При переходе во вкладку “статус принтеров” администратору отображается веб-интерфейс Klipper любого выбранного принтера. В отличие от пользователя, администратор имеет доступ ко всем функциям подключенных принтеров.

Программное взаимодействие модулей системы

На рисунке 3 представлена схема взаимодействия программных модулей многопользовательской системы удаленной 3D печати. Программная часть веб-приложения, расположенного на медиасервере, реализована на языке программирования PHP 7.4.2, внутренняя система обработки заявок и построения очереди печати реализован с помощью языков программирования PHP 7.4.2, Node.js и программной оболочки Bash.

Как видно из рисунка 3, на медиасервере установлена система контейнеризации Docker. В контейнере находится веб-приложение, интерфейс взаимодействия с камерой, который наследует параметры, используемые в основной прошивке устройства, а так

Добавить новый заказ Видеонаблюдение История заказов Статус принтеров

Наблюдение за процессом печати можно выполнять с монитора, расположенного ниже

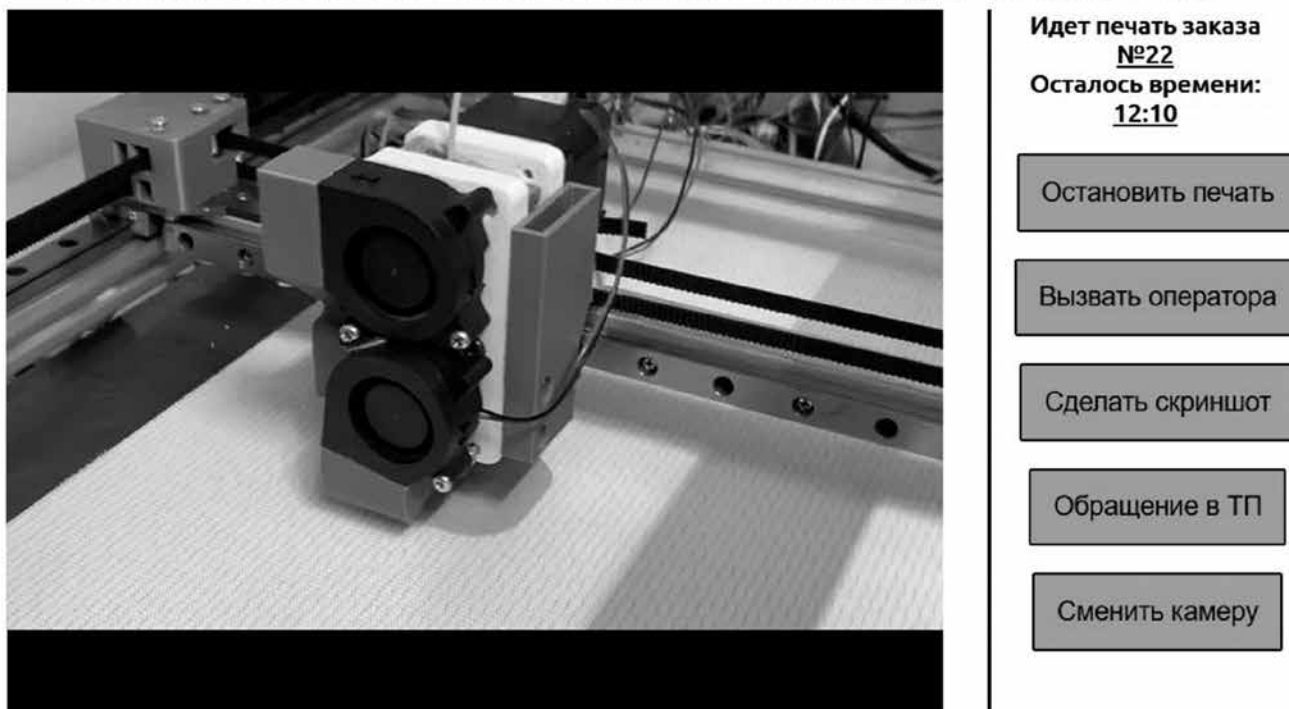


Рис. 2. Окно панели управления системой 3D печати

же хранилище для файлов, которые загружаются на печать. При сборке контейнера программа Docker полностью упаковывает в конечный образ все нужные для работы ПО зависимости и приложения с настройками, которые затрагиваются при функционировании разработанного программного обеспечения [15].

IP-видеокамера, подключенная к медиасерверу по защищенному каналу, передает видеопоток в контейнер, откуда он транслируется пользователю в панель веб-приложения, как показано на рисунке 2.

Рассмотрим структуру взаимодействия пользователя с системой. После авторизации выполняется переход в панель управления. Выбрав файл для загрузки и настроив все нужные пользователю параметры происходит получение запроса на постановку в очередь на печать. Если файл соответствует расширению .g и не превышает объем 400 МБ происходит процесс постановки модели в очередь на печать. В то же время формируется конфигурационный файл, который содержит следующую информацию: номер позиции в очереди, номер пользователя, тип материала, имя g-файла и время, затрачиваемое на печать. Файлы находятся внутри хранилища контейнера на базе библиотеки "vsftpd" с помощью которой можно обеспечивать политику разграниченного доступа к файловой системе вложенных в него элементов. Структура рабочей области библиотеки "vsftpd" состоит из двух папок: "user_files" (папка для хранения g-файлов) и "config" (папка для хранения конфигурационных файлов). Таким образом, после загрузки модели на сервер происходит его ка-

талогизация в очереди, путем создания файлов-маркеров для дальнейшей работы с ними. После постановки модели в очередь пользователю отправляется информация о месте в очереди, рассчитанном времени начала и окончания печати. Эту информацию пользователь может увидеть во вкладке "История заказов".

После постановки файла в очередь на печать медиасервер запускает цикл опросов состояния принтера. Схема взаимодействия платы управления с сервером изображена на рисунке 4. Медиасервер с интервалом 1 с отправляет POST-запрос на плату управления для уточнения статуса принтера. Если принтер занят печатью или идет подготовка к печати, на медиасервер отправляется POST-ответ, содержащий информацию о продолжительности печати и времени окончания печати. Если принтер свободен, или печать только что завершена (модель еще находится в зоне печати), отправляется POST-ответ на разрешение следующей печати. После получения данного ответа, в котором содержится структурный путь нахождения следующего файла-маркера внутри локальной сети, происходит считывание этого файла и, после получения всех нужных конфигурационных данных из него происходит передача g-файла и параметров печати из "vsftpd" хранилища на плату управления процессом печати по протоколу ssh. Таким образом, сразу же после печати предыдущей модели и переноса ее в область охлаждения начинается печать следующей. Это позволяет сэкономить время на разогрев экструдера и термокамеры. Если же во время загрузки файла возникают

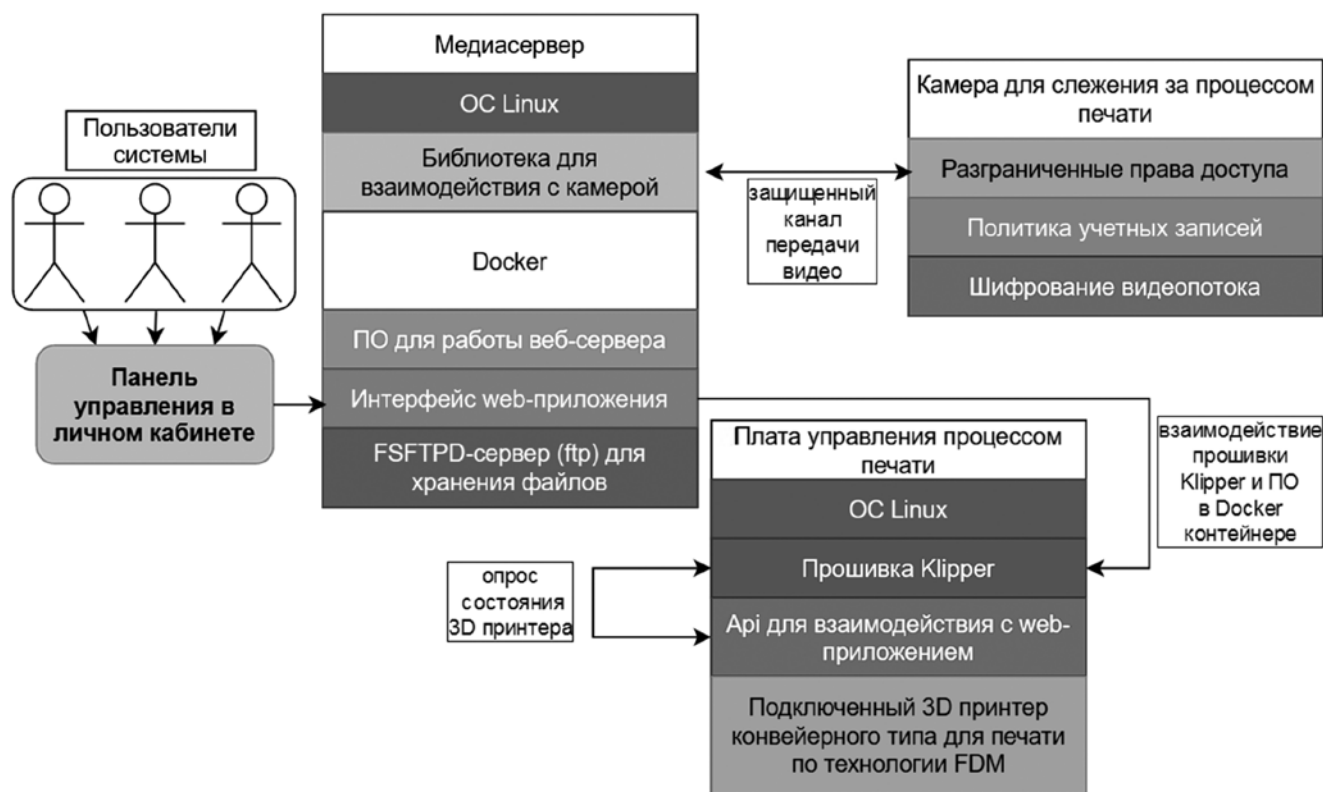


Рис. 3. Схема взаимодействия программных модулей

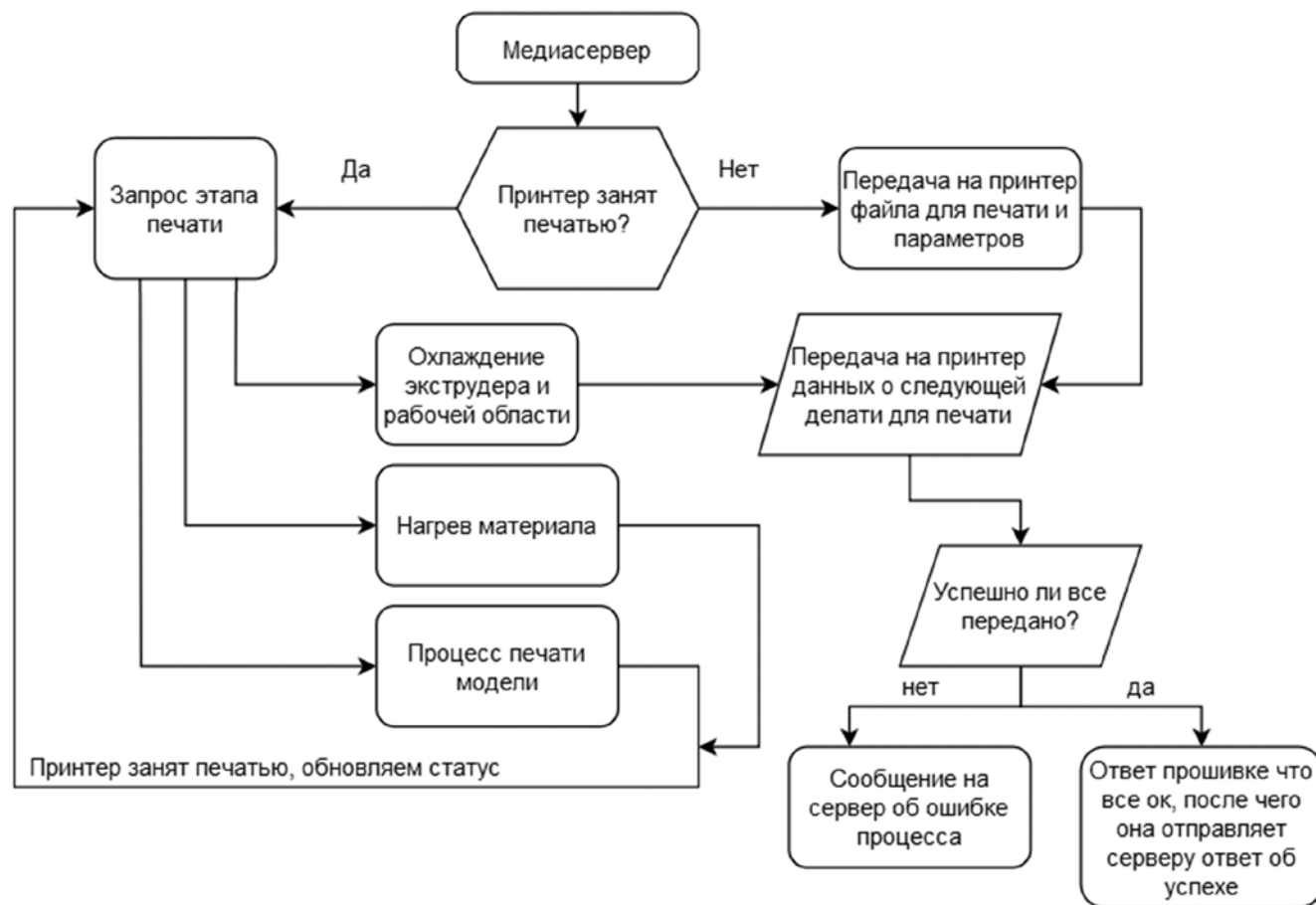


Рис. 4. Схема взаимодействия медиасервера с платой управления принтером

ошибки об этом сразу же сообщается оператору системы. Подключение нескольких принтеров к системе позволяет существенно повысить надежность, так как в случае выхода из строя одного из принтеров и отсутствия POST-ответов от платы управления отправка файла на печать осуществляется исправным принтерам. После окончания печати медиасервер сохраняет изображение готовой модели и оповещает пользователя об окончании печати.

Администратор имеет право пользоваться всеми настройками принтера. Взаимодействие администратора с платой управления осуществляется через защищенный протокол https, доступ предоставляется в соответствующей вкладке “Статус принтеров”.

Заключение

Разработанная многопользовательская система удаленной 3D печати на базе медиасервера позволяет получить доступ к функциям конвейерной 3D печати в любое время, в любой точке планеты при наличии устройства с выходом в интернет. В системе реализованы возможности постановки модели в очередь на печать удаленно, наблюдения за процессом печати и удаленного управления параметрами печати.

Повышенная надежность системы достигается путем пакетного разграничения прав доступа к элементам и множественных проверок в процессе печати.

Для дополнительной защиты доступа данных пользователей была использована система разграничения доступа к данным пользователя через токены аутентификации в виде логина и пароля и изоляции ключевых элементов системы между собой, что привело к хорошей устойчивости системы от разного рода атак и попыток взлома, сводя фиктивные подключения к нулю.

Литература

1. Толкачев, С.А. Сетевая промышленная политика в эпоху новой индустриальной революции / С.А. Толкачев // Журнал НЭА. – 2018. – №. 3. – С. 39.
2. Зленко М.А., Нагайцев М.В., Довбыш В.М. Аддитивные технологии в машиностроении. – 2015.
3. Prendergast, M.E. Recent advances in enabling technologies in 3D printing for precision medicine / M.E. Prendergast, J.A. Burdick // *Advanced Materials*. – 2020. – Vol. 32. – N. 13. – P. 1902516.

4. Karakurt, I. 3D printing technologies: techniques, materials, and post-processing / I. Karakurt, L. Lin // *Current Opinion in Chemical Engineering*. – 2020. – Vol. 28. – P. 134-143.
5. Wang, Y.C. Advanced 3D printing technologies for the aircraft industry: a fuzzy systematic approach for assessing the critical factors / Y.C. Wang, T. Chen, Y.L. Yeh // *The International Journal of Advanced Manufacturing Technology*. – 2019. – Vol. 105. – N. 10. – P. 4059-4069.
6. Kristiawan, R.B. A review on the fused deposition modeling (FDM) 3D printing: Filament processing, materials, and printing parameters / R.B. Kristiawan [et al.] // *Open Engineering*. – 2021. – Vol. 11. – N. 1. – P. 639-649.
7. Yadav, P. Binder jetting 3D printing of titanium aluminides based materials: a feasibility study / P. Yadav [et al.] // *Advanced Engineering Materials*. – 2020. – Vol. 22. – N. 9. – P. 2000408.
8. Karnati, S. On the feasibility of tailoring copper–nickel functionally graded materials fabricated through laser metal deposition / S. Karnati [et al.] // *Metals*. – 2019. – Vol. 9. – N. 3. – P. 287.
9. Wickramasinghe, S. FDM-based 3D printing of polymer and associated composite: A review on mechanical properties, defects and treatments / S. Wickramasinghe, T. Do, P. Tran // *Polymers*. – 2020. – Vol. 12. – N. 7. – P. 1529.
10. Go, J. Rate limits of additive manufacturing by fused filament fabrication and guidelines for high-throughput system design / J. Go [et al.] // *Additive Manufacturing*. – 2017. – Vol. 16. – P. 1-11.
11. 2022 Trends in 3D printing, forecasts from additive manufacturing experts and leaders. [электронный ресурс]. URL: <https://3dprintingindustry.com/news/2022-trends-in-3d-printing-forecasts-from-additive-manufacturing-experts-and-leaders-202426/> (Дата обращения 01.06.22).
12. Thomas, D.S. Costs and cost effectiveness of additive manufacturing / D.S. Thomas [et al.] // *NIST special publication*. – 2014. – Vol. 1176. – P. 12.
13. Blackbelt 3D. [электронный ресурс]. URL: <https://blackbelt-3d.com/> (Дата обращения 01.06.22).
14. Messimer, S.L. Characterization and processing behavior of heated aluminum-polycarbonate composite build plates for the FDM additive manufacturing process / S.L. Messimer [et al.] // *Journal of Manufacturing and Materials Processing*. – 2018. – Vol. 2. – N. 1. – P. 12.
15. Poulton, N. Docker deep dive. – JJNP Consulting Limited, 2019. – 226 p.

Поступила в редакцию 9.06.2022 г.

	НАИМЕНОВАНИЕ ТОВАРА	НАЗВАНИЕ КОМПАНИИ, АДРЕС, ТЕЛЕФОН
КВАРЦЕВЫЕ РЕЗОНАТОРЫ, ГЕНЕРАТОРЫ, ФИЛЬТРЫ, ПЬЕЗОКЕРАМИЧЕСКИЕ И ПАВ ИЗДЕЛИЯ		
1.1	Любые кварцевые резонаторы, генераторы, фильтры (отечественные и импортные)	 ·ALNAR· УП «Алнар» +375 (17) 227-69-97 +375 (17) 227-28-10 +375 (17) 227-28-11 +375 (29) 644-44-09 alnar@tut.by www.alnar.net
1.2	Кварцевые резонаторы Jauch под установку в отверстия и SMD-монтаж	
1.3	Кварцевые генераторы Jauch под установку в отверстия и SMD-монтаж	
1.4	Термокомпенсированные кварцевые генераторы	
1.5	Резонаторы и фильтры на ПАВ	
1.6	Пьезокерамические резонаторы, фильтры, звонки, сирены	

УНП 100191870

СПЕЦПРЕДЛОЖЕНИЕ		
2.1	Большой выбор электронных компонентов со склада и под заказ. Микросхемы производства Xilinx, Samsung, Maxim, Atmel, Altera, Infineon и пр. Термоусаживаемая трубка, диоды, резисторы, конденсаторы, паяльная паста, кварцевые резонаторы и генераторы, разъемы, коммутация и др.	 ПОСТАВКА ЭЛЕКТРОННЫХ КОМПОНЕНТОВ ЧТУП «Чип электроникс» +375 (17) 269-92-36 chipelectronics@mail.ru www.chipelectronics.by
2.2	Широчайший выбор электронных компонентов (микросхемы, диоды, тиристоры, конденсаторы, резисторы, разъемы в ассортименте и др.)	Группа компаний «Альфа-лидер» +375 (17) 391-02-22 +375 (17) 391-03-33. www.alider.by

УНП 191142740

УНП 192321381

ЭЛЕКТРОННАЯ И ЭЛЕКТРОТЕХНИЧЕСКАЯ ПРОДУКЦИЯ		
3.1	Комплексная поставка электронных компонентов	 ТУП «Альфа-чип Лимитед» +375 (17) 366-76-16 analog@alfa-chip.com www.alfa-chip.com
3.2	Датчики, сенсоры и средства автоматизации	
3.3	Светодиодные индикаторы, TFT, OLED и ЖК-дисплеи и компоненты для светодиодного освещения	
3.7	Мощные светодиоды (EMITTER, STAR), сборки и модули мощных светодиодов, линзы ARLIGHT	 ООО «СветЛед решения» +375 (17) 214-73-27 +375 (17) 214-73-55 info@belaist.by www.belaist.by
3.8	Управление светом: RGB-контроллеры, усилители, диммеры и декодеры	
3.9	Источники тока AC/DC для мощных светодиодов (350/700/ 100-1400 mA) мощностью от 1 W до 100 W ARLIGHT	
3.10	Источники тока DC/DC для мощных светодиодов (вход 12-24V) ARLIGHT	
3.11	Источники напряжения AC/DC (5-12-24-48 V от 5 до 300 W) в металлическом кожухе, пластиковом, герметичном корпусе ARLIGHT, HAITAIK	
3.12	Светодиодные ленты, линейки открытые и герметичные, ленты бокового свечения, светодиоды выводные ARLIGHT	
3.13	Светодиодные лампы E27, E14, GU 5.3, GU 10 и др.	
3.14	Светодиодные светильники, прожекторы, алюминиевый профиль для светодиодных изделий	

УНП 192525135

УНП 191672332

СПЕЦПРЕДЛОЖЕНИЕ		
3.15	Поставка со склада и под заказ: микросхемы TEXAS INSTRUMENTS, INTERSIL, EM Marin, FREESCALE, XILINX, ALTERA, CHINFА, реле GRUNER, кварцевые резонаторы KDS, MICRO KRISTAL, батарейки и аккумуляторы, держатели RENATA, XENO, PKCELL, модемы HUAWEI, QUECTEL, системы на модуле (одноплатные компьютеры) отладки, беспроводные модули SECO, INMIS, SMK, SAURIS, TORADEX, накопители на флэш-памяти INNODISK, герконы COMUS, COTO, разъемы KEYSTONE, HIROSE и др. Техническая поддержка, поставка бесплатных образцов, проектные цены.	 ООО «БелСКАНТИ» +375 (17) 256-08-67, +375 (17) 398-21-62 nab@scanti.ru www.scanti.com

УНП 190813939

3.16	Индуктивные, емкостные, оптоэлектронные, магнитные, ультразвуковые, механические датчики фирмы Balluff (Германия)
3.17	Блоки питания, датчики давления, разъемы, промышленная идентификация RFID, комплектующие фирмы Balluff (Германия)
3.18	Магнитострикционные, индуктивные, магнитные измерители пути, лазерные дальномеры, индуктивные сенсоры с аналоговым выходом, инклинометры фирмы Balluff (Германия)
3.19	Инкрементальные, абсолютные, круговые магнитные энкодеры фирмы Lika Electronic (Италия)
3.20	Абсолютные и инкрементальные магнитные измерители пути, УЦИ (устройство цифровой индикации), тросиковые блоки, муфты, угловые актуаторы фирмы Lika Electronic (Италия)
3.21	Автоматические выключатели, УЗО, дифавтоматы, УЗИП, выключатели нагрузки фирмы Schneider Electric (Франция)
3.22	Контакты, промежуточные реле, тепловые реле перегрузки, реле защиты, автоматические выключатели защиты двигателя фирмы Schneider Electric (Франция)
3.23	Кнопки, переключатели, сигнальные лампы, посты управления, джойстики, выключатели безопасности, источники питания, световые колонны фирмы Schneider Electric (Франция)
3.24	Универсальные шкафы, автоматические выключатели, устройства управления и сигнализации, УЗО и дифавтоматы, промежуточные реле, выключатели нагрузки, контакторы, предохранители, реле фирмы DEKraft

АВТОМАТИКА
Ц · Е · Н · Т · Р

ООО «Автоматика центр»

+375 (17) 218-17-98

+375 (17) 218-17-13

sos@electric.by

www.electric.by

УНП 191087188

Беззеркальный фотоаппарат Olympus E-M10 Mark IV Kit 14-150 mm

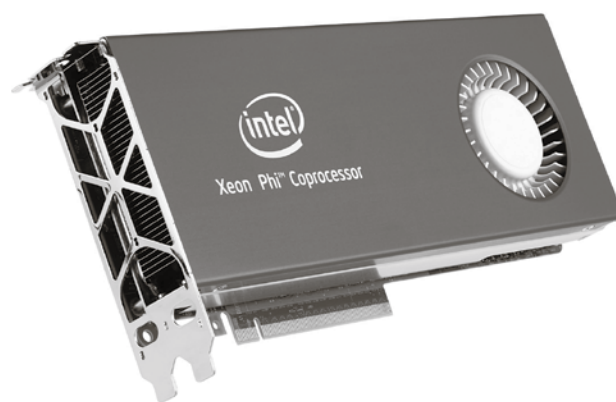


Компактная беззеркалка для туризма и влога...

OLYMPUS
WWW.OLYMPUS.COM

Сократите свой путь к открытию!

Intel Xeon Phi сопроцессор



ПОДРОБНОСТИ
НА INTEL.COM

**XXVI БЕЛОРУССКИЙ
ЭНЕРГЕТИЧЕСКИЙ и
ЭКОЛОГИЧЕСКИЙ
ФОРУМ**

energyexpo.by

ENERGY EXPO

ЭНЕРГЕТИКА
ЭКОЛОГИЯ
ЭНЕРГОСБЕРЕЖЕНИЕ
ЭЛЕКТРО

green
industry

ИННОВАЦИОННЫЕ
ПРОМЫШЛЕННЫЕ
ТЕХНОЛОГИИ

etrans

САЛОН
ИННОВАЦИОННОГО
ТРАНСПОРТА

11-14.10.2022

Минск, пр. Победителей, 20/2

А л в ф а Ч И П Л И М И Т Е Д

*Новые возможности
ваших идей*

- Электронные компоненты
- Средства автоматизации
- Датчики, сенсоры
- Светодиодные индикаторы, TFT, OLED и ЖКИ дисплеи
- Компоненты для светодиодного освещения

Прямые поставки
от мировых производителей

Разработка и техническая
поддержка новых проектов



220012, г. Минск, ул. Сурганова, 5а, 1-й этаж
Тел./факс: +375 17 366 76 01, +375 17 366 76 16
www.alfa-chip.com
www.alfacomponent.com