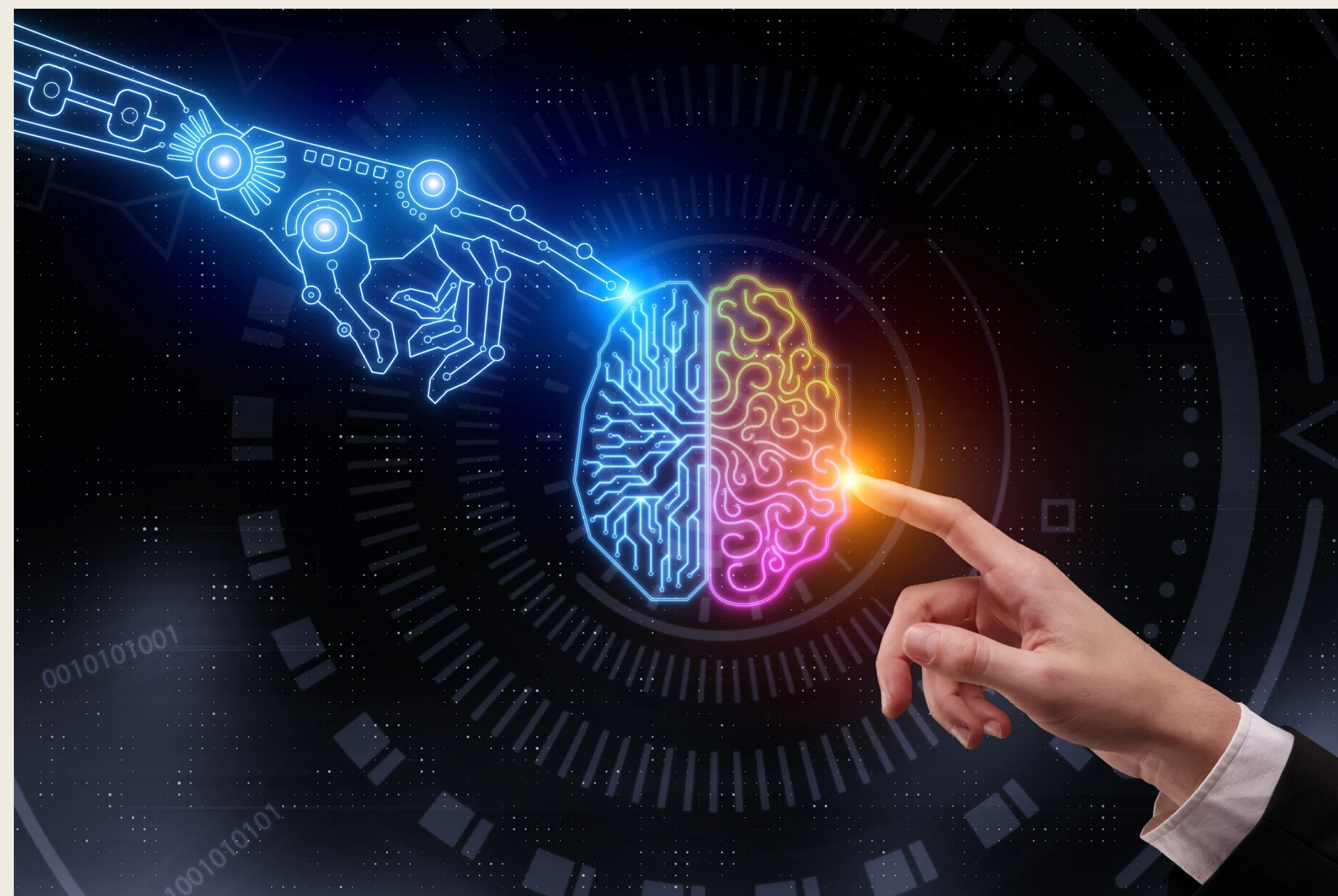




МФК

Правовое регулирование искусственного интеллекта на международном уровне

Татьяна Григорьевна Левченко,
кандидат юридических наук, доцент



ИИ

Проблема определения ИИ

- моделирование поведения человека
- «Мы создаем инструменты и формируем себя посредством их использования» (Дж. Янг)

Основа основ - Н. Винер

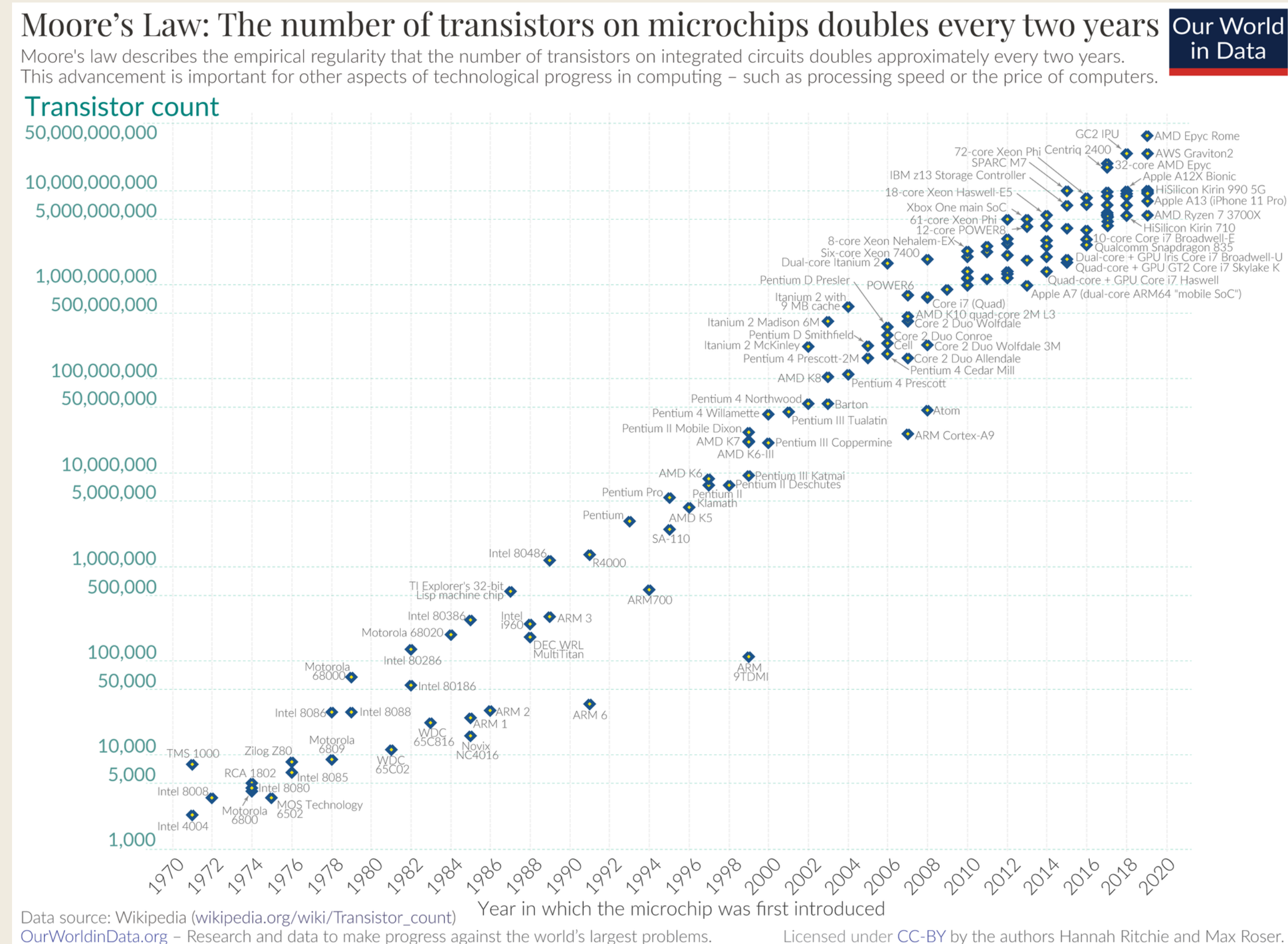
- кибернетика Н. Винера - представление об автоматическом и саморегулируемом управлении;
 - возникновение нового фашизма, который опирается на *machine gouverner*;
 - недостаток внимания к когнитивному элементу
-

Основа основ - Н. Винер

- сложные системы состоят из взаимосвязанных контуров обратной связи, где обмен сигналами между подсистемами порождает комплексное, но стабильное поведение;
 - информация - центральное звено управления поведением сложных систем;
-

Ошибки Винера

- недооценка потенциала цифровых вычислений;
- закон Мура - удвоение мощности процессоров каждые два года (с 1950 г.)
- предел закона Мура



Закон Мура

- предел закона Мура обусловлен пределами, налагаемыми физикой;
- выбор между эффективностью, скоростью и рациональностью

**Немного
передвинул
картинку
влево**



**Весь текст съехал, 4
новые страницы
открылись, орбита Земли
сместилась на 2 метра,
где-то вдалеке зазвучали
сирены**

Основа основ

- кибернетика наблюдаемые систем ➡ кибернетика систем наблюдения (Х. фон Ферстер);
 - коммуникация - все процедуры и способы, с помощью которых один разум может влиять на другой (К. Шенон);
-

Новый подход к ИИ

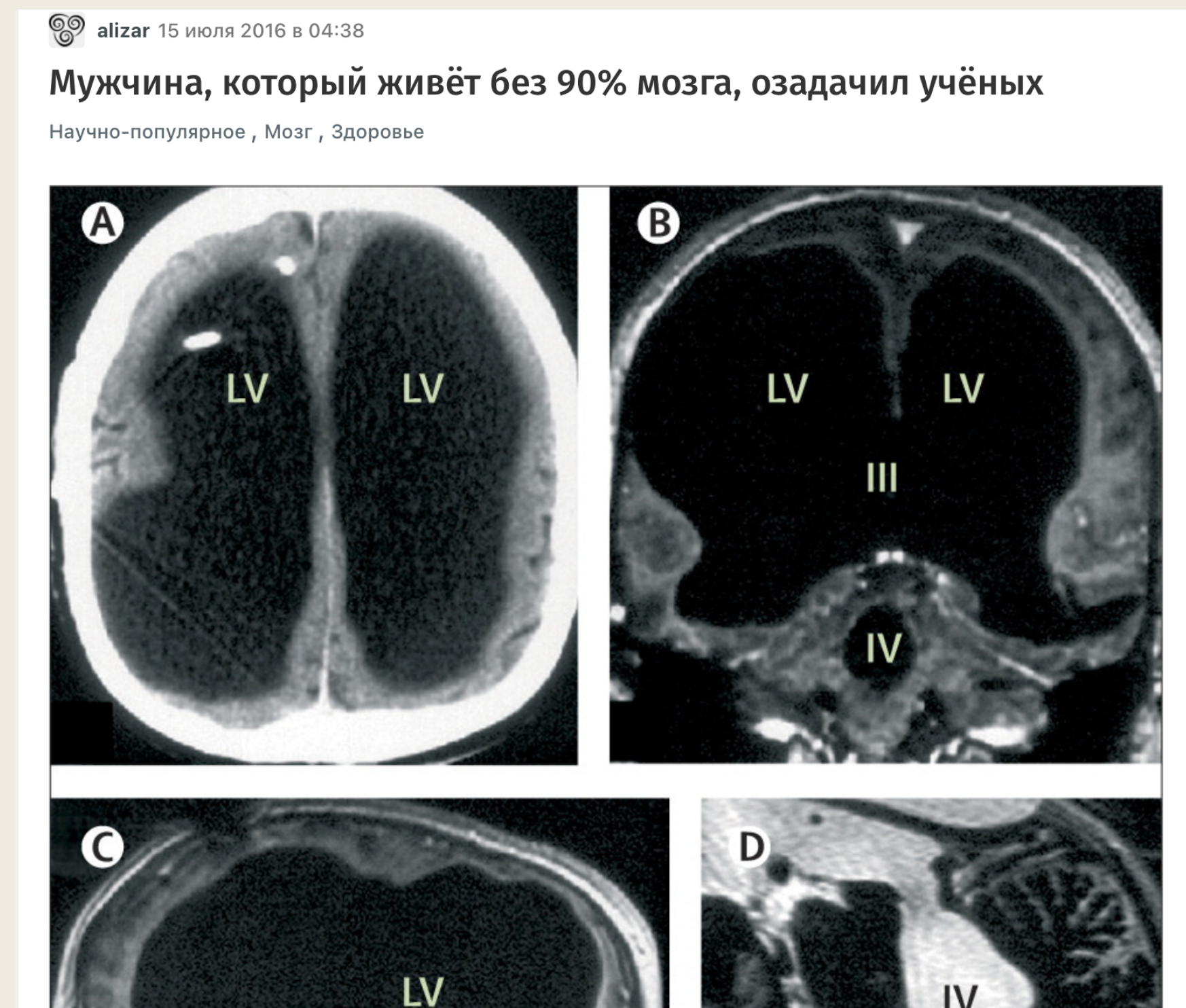
- Дж. Перл - вероятностная модель машинного мышления - «пересматривать свои убеждения в свете новых свидетельств»;
 - ИИ должен имитировать человеческое мышление и компетентность;
-

Человеческое сознание

- ментальные репрезентации;
 - выбор информации для глобального доступа всего организма;
 - самоконтроль.
-

Теория биологической сущности сознания

- теория клауструм;
- теория зрительной коры;
- [https://www.thelancet.com/journals/lancet/article/PIIS0140-6736\(07\)61127-1/fulltext](https://www.thelancet.com/journals/lancet/article/PIIS0140-6736(07)61127-1/fulltext)

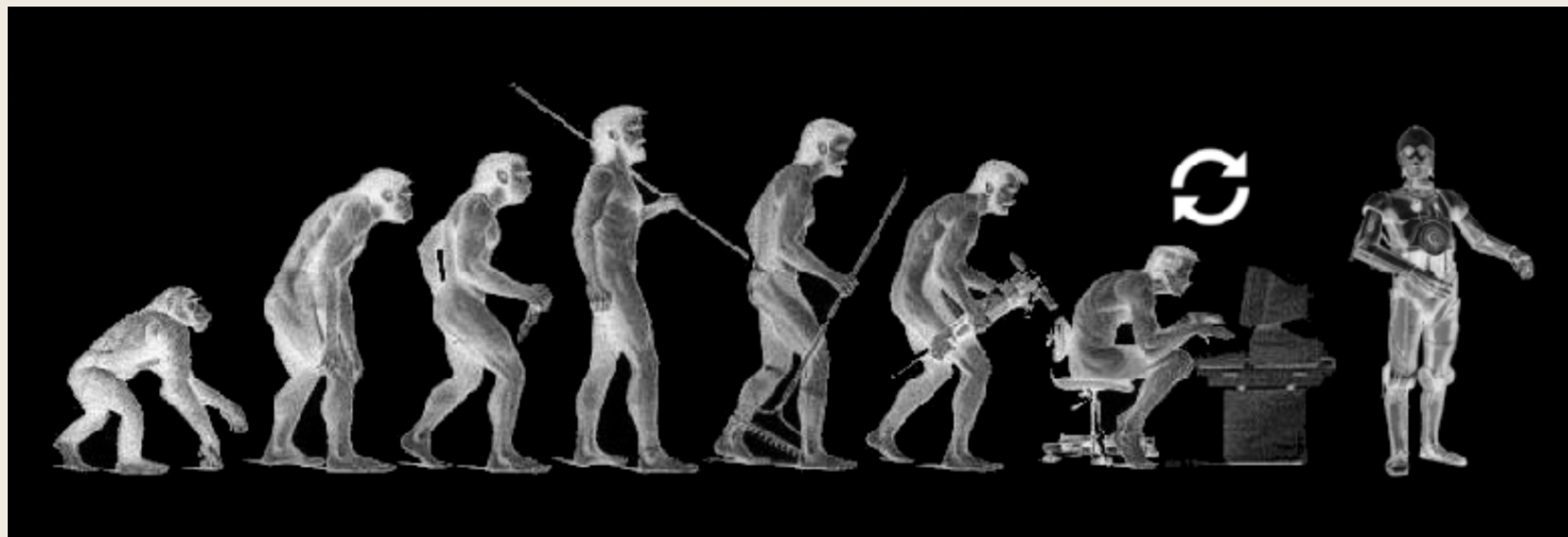


Опасности ИИ

- сможем ли мы контролировать самосовершенствующийся ИИ?
 - «опасность не в машинах, которые становятся все больше похожими на людей, а в людях, с которыми обращаются как с машинами» (Н. Винер);
 - автоматизация производства приводит не только к повышению производительности труда, но и к безработице.
-

Опасность ИИ

- технологическая сингулярность - момент, когда технологическое развитие станет неуправляемым и необратимым;



Опасности ИИ

- проблемы технологического прогнозирования;
- достижимость сверхчеловеческого уровня для ИИ;

Проблемы постановок целей для ИИ

- машины должны быть в неведении истинных целей программистов? (Р. Стюарт)
 - может ли человечество уступить машинам право распоряжаться своей судьбой? (Н. Винер)
 - разумная машина с выключателем будет стремиться деактивировать выключатель (С. Омохундро)
-

Зачет

проект международной конвенции «Об общих принципах регулирования искусственного интеллекта»

5 человек в команде (желательно с разных факультетов)

до 1 декабря на почту levchenko2011@yandex.ru

в теме письма - МФК; название файла - фамилии;

антиплагиат
