

§ 1–2. Искусственный интеллект

Искусственный интеллект (ИИ) — это способность машин выполнять задачи, которые раньше считались возможными только для человека: рассуждать, принимать решения, учиться и даже творить.

Термин впервые предложил в 1956 году американский учёный Джон Маккарти. В английском выражении *artificial intelligence* слово *intelligence* скорее означает «разумное мышление», чем привычное нам «интеллект».

Ещё раньше, в 1950 году, выдающийся учёный Алан Тьюринг задал провокационный вопрос: «*Может ли машина мыслить?*». В своей статье он описал процедуру, с помощью которой можно проверить, сравнялась ли машина с человеком по уровню рассуждений. Этот метод мы теперь знаем как **тест Тьюринга**.

ИИ — это не магия, а результат работы программ, управляющих машинами. Главная цель таких систем — **имитировать человеческое поведение и способность рассуждать**.

Почему для ИИ важен человеческий мозг

Многие исследователи считают, что, поняв, как работает мозг, мы сможем построить полноценный искусственный интеллект.

Ведь если удастся воспроизвести процессы обучения, мышления и принятия решений, то можно будет создать машину, которая:

- учится на своём опыте,
- принимает решения,
- решает творческие задачи.

Такие системы уже применяются: от «умных» помощников и переводчиков до автомобилей, которые сами ездят по дорогам.

Машинное обучение – сердце ИИ

Одним из ключевых направлений является **машинное обучение**. Оно позволяет компьютерам учиться на больших массивах данных.

Благодаря этому технологии способны:

- распознавать лица и речь,
- переводить тексты,
- находить закономерности в данных,
- делать прогнозы.

По сути, машина не просто выполняет заложенные команды, а сама учится на примерах — как ребёнок, который на опыте узнаёт, что такое «собака» или «стол».

Нейронные сети – мозг внутри компьютера

Особая разновидность машинного обучения — это **искусственные нейронные сети (ИНС)**.

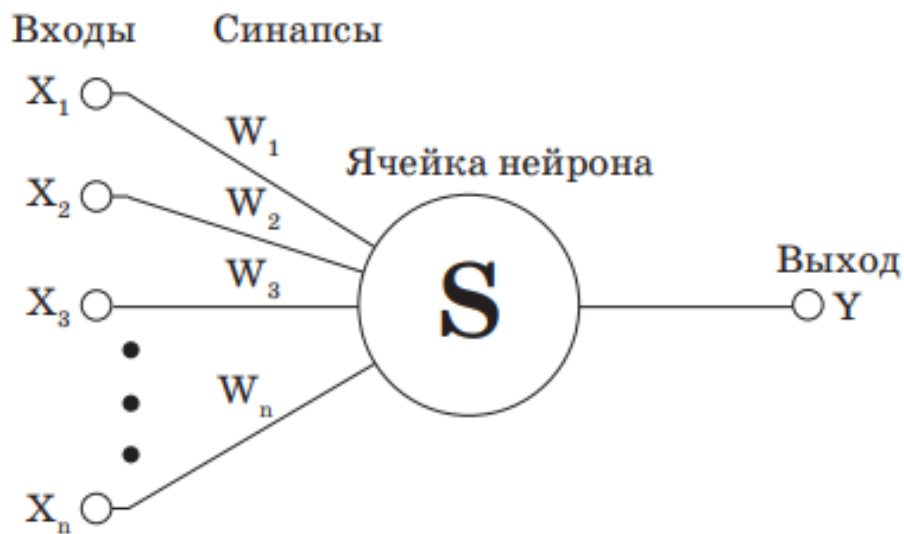


Схема 1. Модель ИНС

Что мы видим на схеме:

1. Входы ($X_1, X_2, \dots, X_n$)

Это входные данные, которые получает нейрон.

Пример: если сеть обучается распознавать изображение кошки, то входами могут быть яркости отдельных пикселей.

2. Синапсы ($W_1, W_2, \dots, W_n$)

Это «веса» связей между входами и нейроном.

Каждый вес показывает, насколько важен соответствующий вход.

- Если вес большой → вход сильно влияет на результат.
- Если маленький → влияние слабое.
- Может быть даже отрицательным → вход работает «в обратную сторону».

3. Ячейка нейрона (S)

Здесь происходит **обработка данных**.

Нейрон суммирует все входные данные, умноженные на их веса:

4. Выход (Y)

Это результат работы нейрона.