

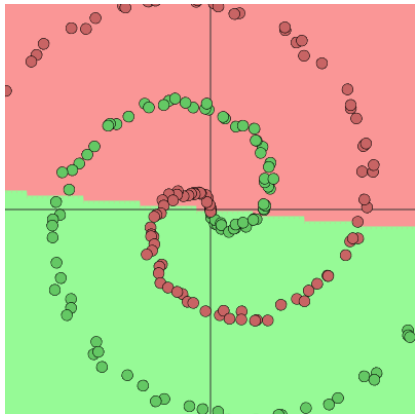
# Искусственные нейронные сети для обработки текстов

Трифонов Владислав

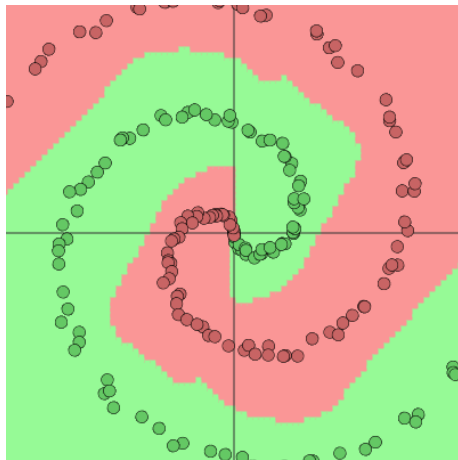
25 сентября 2019 г.

Нейронные сети – мощные алгоритмы машинного обучения, широко используемые при решении задач автоматической обработки текстов, изображений, видео и звуков.

Они могут находить сложные зависимости в данных, что недоступно для традиционных методов ML.

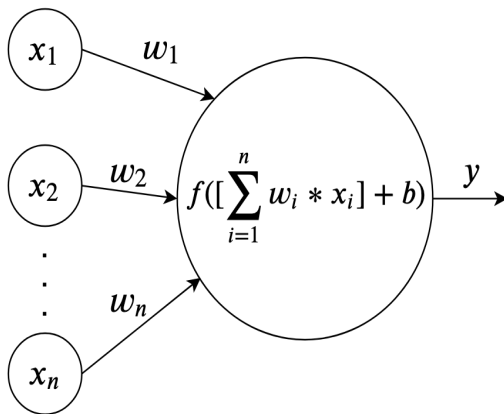


Линейный классификатор



Нейронная сеть с tanh скрытым слоем

# Модель нейрона



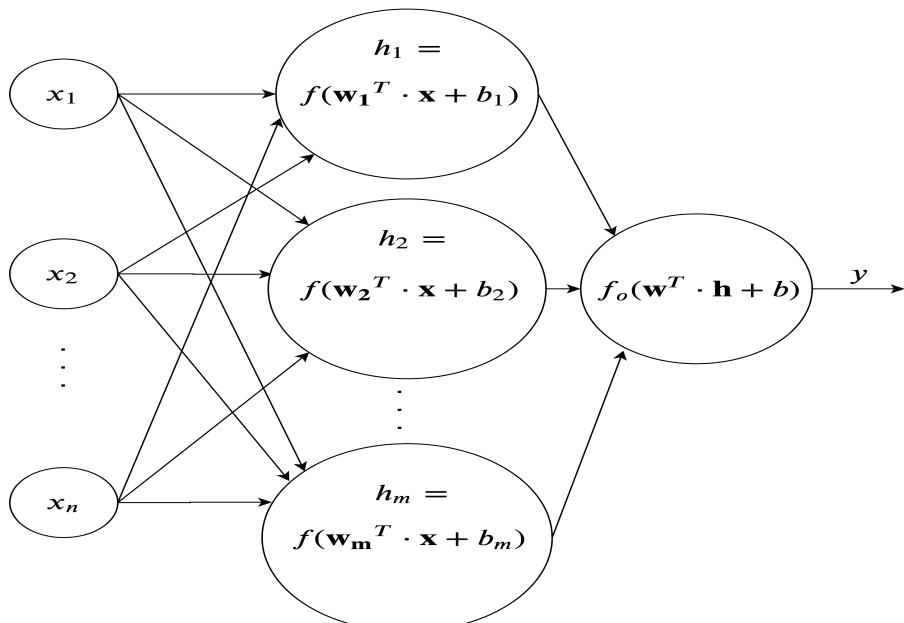
- W. McCulloch, W. Pitts – 1943
- $y = f(\mathbf{w}^T \cdot \mathbf{x} + b)$
- $\mathbf{w}$  - веса связей
- $b$  - смещение (bias)
- $f$  - функция активации

- Нейрон линейно комбинирует входные признаки и с помощью функции активации генерирует число – сигнал, который можно использовать как ответ классификатора или как новый признак
- Линейные модели (линейная регрессия, линейный классификатор, логистическая регрессия) – простейшие нейронные сети с одним нейроном

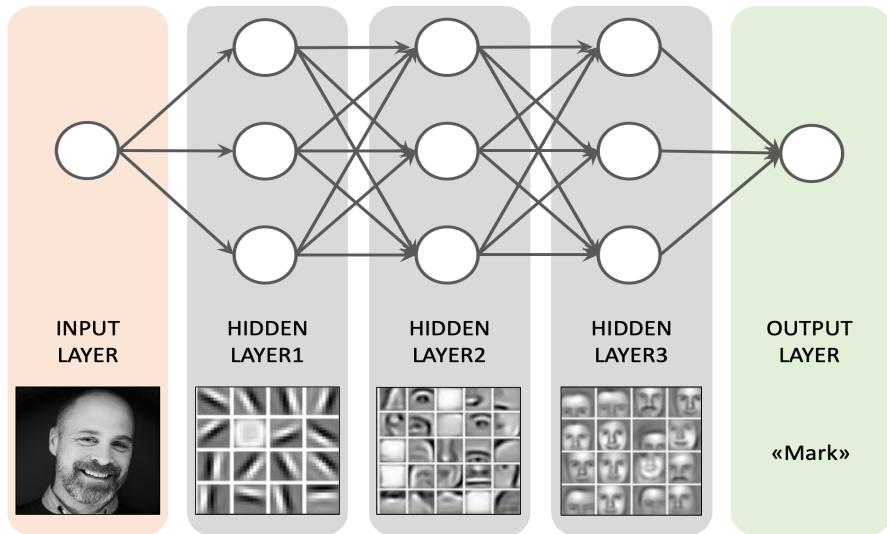
## Идея

Разработать модель, которая научиться из исходных признаков извлекать новые и на их основе принимать итоговое решение

# Персептрон с одним скрытым слоем



# Персептрон



<sup>0</sup><https://www.bouvet.no/bouvet-deler/an-introduction-to-deep-learning>

- Скрытый слой можно записать как  $\mathbf{h} = f(\mathbf{W}\mathbf{x} + \mathbf{b})$ ;  $\mathbf{W} \in \mathbb{R}^{m \times n}$ ;  $\mathbf{x} \in \mathbb{R}^n$ ;  $\mathbf{b}, \mathbf{h} \in \mathbb{R}^m$
- На практике используется не более 2-3 скрытых слоев. Больше параметров – выше вероятность переобучения



$y = f_o(\mathbf{w}^T \cdot \mathbf{h})$  – выходной слой. В качестве  $f_o$  обычно используется функция для принятия решения из линейных моделей:

- $y = \mathbf{w}^T \cdot \mathbf{h}$  – регрессия

- $y = \sigma(\mathbf{w}^T \cdot \mathbf{h}) = \frac{1}{1 + \exp -(\mathbf{w}^T \cdot \mathbf{h})}$  – бинарная классификация

- $\mathbf{y} = \text{softmax}(\mathbf{W}\mathbf{h})$ ,  $\mathbf{W} \in \mathbb{R}^{c \times m}$  – классификация на  $c > 2$  классов

$$\text{softmax}_i(\mathbf{x}) = \frac{\exp(x_i)}{\sum_{k=1}^{|\mathbf{x}|} \exp(x_k)}$$

- Пусть  $(x_i, \tilde{y}_i), i \in [1, n]$  – обучающая выборка, где  $x_i$  – объекты,  $\tilde{y}_i$  – правильные ответы
- Пусть  $y = y_\theta(x)$  – дифференцируемая нейронная сеть с параметрами  $\theta$ , которые мы хотим обучить
- Параметры  $\theta$  должны быть такими, чтобы нейронная сеть выдавала правильные ответы на элементах обучающей выборки

- Для оценки погрешности алгоритма выбирается функция ошибки  $l = l(y_i, \tilde{y}_i)$ ,  $l \in \mathbb{R}$ . Из ранее рассмотренных – MSE, cross-entropy, hinge
- По всей обучающей выборке строим функционал качества, применяя нейронную сеть и сравнивая ее ответы с правильными:

$$L = L(\theta) = \frac{1}{n} \sum_{i=1}^n l(y_{\theta}(x_i), \tilde{y}_i) = \frac{1}{n} \sum_{i=1}^n l_i$$

- Т.к.  $L$  – дифференцируемая функция, то можем находить ее минимум по параметрам  $\theta$  с помощью градиентного спуска:  
$$\theta_{i+1} = \theta_i - \alpha \cdot \nabla L(\theta_i)$$

- Вычисление  $L$  при больших  $n$  – ресурсоемкая задача
- Stochastic gradient descent (SGD) – на каждом шаге градиентного спуска случайно выбираем **один** пример из обучающей выборки и обновляем параметры:  $\theta_{i+1} = \theta_i - \alpha \cdot \nabla l_j(\theta_i)$
- SGD – подвержен “выбросам”, поэтому шаг  $\alpha$  для сходимости следует выбирать маленьким, а следовательно процесс обучения может проходить долго
- Вычисление не параллелится по элементам обучающей выборки

- miniBatch SGD – на каждом шаге градиентного спуска случайно выбираем  $m$  примеров ( $1 < m \ll n$ ) из обучающей выборки и обновляем параметры:  $\theta_{i+1} = \theta_i - \alpha \cdot \nabla[\frac{1}{m} \sum_{j=1}^m l_j(\theta_i)]$
- Более устойчив к “выбросам”, поэтому можем выбирать больший шаг  $\alpha$
- Потерю  $l_j(\theta_i)$  можно считать параллельно по  $j$

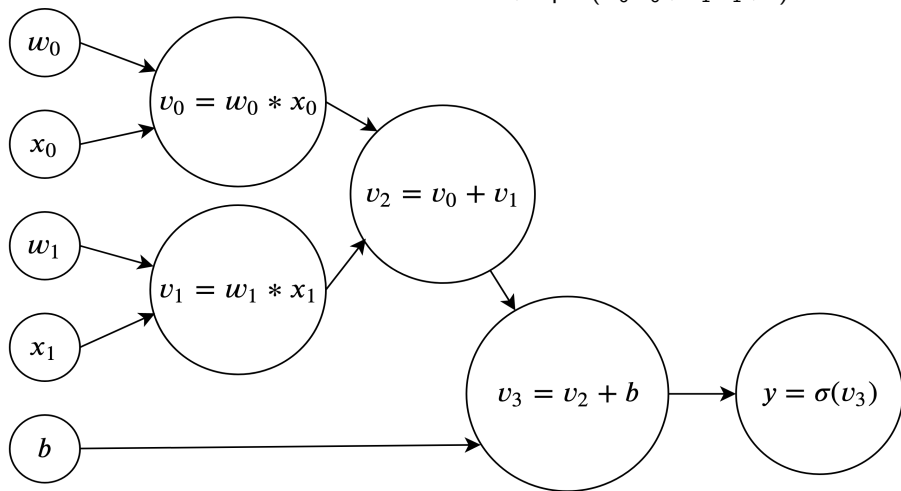
Обучение проводят некоторое количество “эпох” – полных итераций по обучающей выборке и останавливаются, когда качество на валидационной выборке перестает улучшаться

В настоящее время в основном используются адаптивные методы градиентного спуска, т.е. значение  $\alpha$  на каждом шаге меняется в зависимости от предыдущих значений градиента – Momentum, RMSProp, Adam и т.д.

## Вопрос

Как вычислять градиент по каждому параметру сложной нейронной сети (в общем случае, сложной дифференцируемой функции)?

Логистическая регрессия:  $y = \frac{1}{1 + \exp(-(w_0 \cdot x_0 + w_1 \cdot x_1 + b))}$



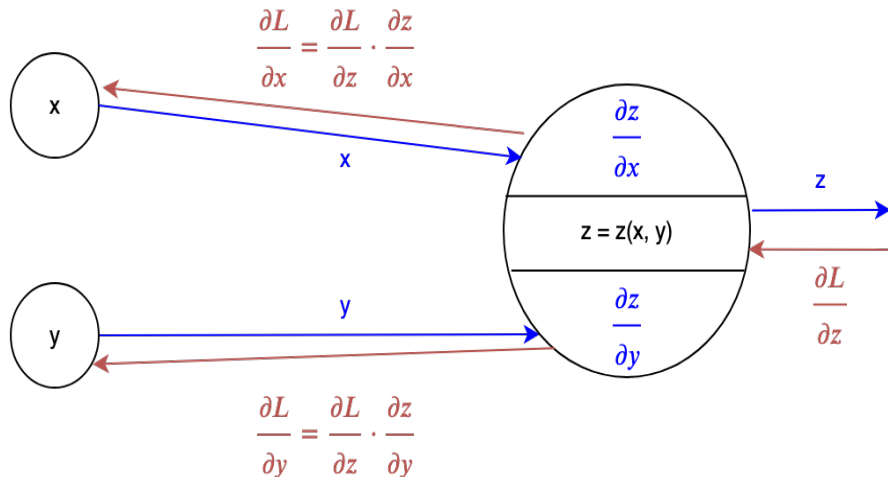
## Backpropagation

На входе граф вычислений некоторой дифференцируемой функции  $y$ .  
Алгоритм:

- 1 Вычисляем значения и частные производные для каждого из узлов (forward pass)
- 2 Вычисляем частные производные реализуемой функции  $y$  по каждому из ее аргументов – узлов графа (backward pass)



# Backpropagation

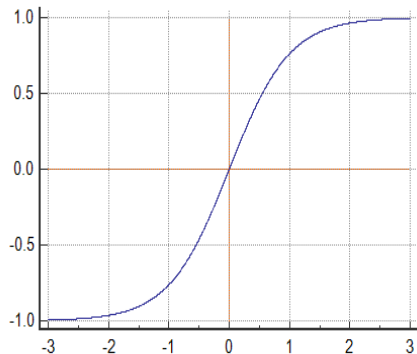


# Проблема затухающих градиентов

При backpropagation в глубокой нейронной сети градиент на первых слоях может быть очень маленьким из-за последовательного умножения на  $|\frac{\partial z}{\partial y}| \ll 1$ , поэтому нейронная сеть не сможет обучиться.

Чтобы с этим бороться, выбирают специальные функции активации и разрабатывают “хитрые” архитектуры.

# Функции активации – tanh

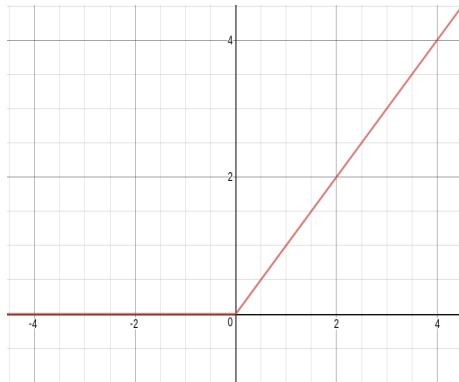


- $y = \tanh(x)$

- $y \in [-1, 1]$

- $\frac{\partial y}{\partial x} \leq 1.0$

# Функции активации – ReLU



- $y = \text{ReLU}(x) = \max(0, x)$

- $y \in [0, +\infty)$

- $\frac{\partial y}{\partial x} \in \{0, 1\}$

- эффективное  
вычисление

- Широко используется в  
скрытых слоях

Перед обучением нейронной сети нужно проинициализировать ее параметры  $\theta$ . Обычно используется случайная инициализация небольшими числами из интервала  $(-1, 1)$ .

Почему все веса нельзя инициализировать 0? Константой?

# Модельная задача

Будем классифицировать СМС на два класса: **не спам**, **спам**

Nah I don't think he goes to usf, he lives around here though – **не спам**

You are a winner U have been specially selected 2 receive £1000 or a 4\* holiday (flights inc) speak to a live operator 2 claim 0871277810910p/min (18+) – **спам**

## Постановка задачи

Построить признаковое пространство  $X$  и алгоритм  $y : X \rightarrow \{0, 1\}$

Прежде, чем приступать к решению задачи, нужно понять, как оценивать качество разработанного метода.

## Важно

Метрика должна отражать реальное качество, а не быть “высокой цифрой”, даже если задача решена плохо

# Метрики качества бинарной классификации

## Типы ошибок

- **tp** – true positive, метод – 1, gold метка – 1
- **tn** – true negative, метод – 0, gold метка – 0
- **fp** – false positive, метод – 1, gold метка – 0
- **fn** – false negative, метод – 0, gold метка – 1

- $accuracy = \frac{tp + tn}{tp + tn + fp + fn}$  – плоха при несбалансированности классов

- $precision = \frac{tp}{tp + fp}$ ,  $recall = \frac{tp}{tp + fn}$ ,  $F1 = 2 * \frac{precision * recall}{precision + recall}$

- Можно AUC-ROC<sup>a</sup>

---

<sup>a</sup><https://developers.google.com/machine-learning/crash-course/classification/roc-and-auc>



- Провести токенизацию – разбить текст на словоформы (токены)
- Построить признаковый вектор для текста, как последовательности токенов

- TF-IDF (TF — term frequency, IDF — inverse document frequency)
- $tf\text{-}idf(t, d, D) = tf(t, d) \cdot idf(t, D)$ , где  $t$  — токен,  $d$  — текущий документ,  $D$  — обучающая выборка
- $tf(t, d) = \frac{n_t}{\sum_k n_k}$  — частота токена в данном документе
- $idf(t, D) = \log \frac{|D|}{|\{d_i \in D | t \in d_i\}|}$

- По обучающей выборке  $D$  строим словарь токенов  $V$  и считаем  $idf(t, D)$  для каждого  $t \in V$
- Теперь можем получить векторное представление  $\mathbf{x}$  для произвольного текста  $d$ :  $\mathbf{x} \in \mathbb{R}^{|V|}$ ,  $x_i = \text{tf-idf}(t_i, d, D)$ ,  $i \in [1, |V|]$
- TF-IDF дает высокий вес токенам, часто встречающимся в конкретном документе, но обладающим низкой частотой относительно всей коллекции
- Таким образом, не обладающие смыслом слова (знаки препинания, соединительные союзы и т.д.) будут иметь низкий вес

- `sklearn.neural_network` – недостаточная конфигурация
- `tensorflow` – высокий порог вхождения
- `keras`, `pytorch`

# Однослойный перцептрон на pytorch

```
import torch

class BinaryClassificationPerceptron(torch.nn.Module):
    def __init__(self, input_size: int, hidden_size: int):
        super().__init__()

        self._hidden = torch.nn.Linear(
            input_size, hidden_size)
        self._activation = torch.nn.ReLU()

        self._out_layer = torch.nn.Linear(hidden_size, 1)
        self._out_sigmoid = torch.nn.Sigmoid()

    def forward(self, sample):
        sample = self._activation(self._hidden(sample))
        return self._out_sigmoid(self._out_layer(sample))
```

Усреднение по 5 запускам с разными случайными инициализациями весов:

- $accuracy = 0.955$
- $precision = 0.889$
- $recall = 0.766$
- $f1 = 0.822$

Статистика одного из запусков:

- $tp = 84$
- $tn = 717$
- $fp = 7$
- $fn = 28$
- $tp + fn = 112$  – положительные примеры
- $tn + fp = 724$  – отрицательные примеры

Исходные тексты содержат много незначимых цифр и большое число ссылок. Попробуем нормализовать текст: заменить все цифры на 0, все ссылки на `<href>`, чтобы конечное признаковое пространство не было разреженным и одинаковые по сути словоформы были действительно одинаковыми.

Усреднение по 5 запускам с разными случайными инициализациями весов:

- $accuracy = 0.968 > 0.955$
- $precision = 0.905 > 0.889$
- $recall = 0.848 > 0.766$
- $f1 = 0.875 > 0.822$

Небольшая предобработка значительно улучшила качество метода



- [https://github.com/vladosdudos/simple\\_spam\\_detection](https://github.com/vladosdudos/simple_spam_detection)
- или  
[https://nbviewer.jupyter.org/github/vladosdudos/simple\\_spam\\_detection/blob/master/simple\\_nns.ipynb](https://nbviewer.jupyter.org/github/vladosdudos/simple_spam_detection/blob/master/simple_nns.ipynb)

[Александр Пушкин | ЛИЧ] родился в [Москве | ЛОК]

IO-кодирование:

I-ЛИЧ I-ЛИЧ O O I-ЛОК

Чтобы подать каждое слово в нейронную сеть, нужно его представить вектором

Пусть  $V$  – перенумерованное множество всех словоформ обучающей выборки (словарь).

Тогда каждому слову можно сопоставить вектор  $\mathbf{x} \in \mathbb{R}^{|V|+1}$ , где  $x_i = 1$ , если это  $i$ -е слово из словаря,  $x_j = 0, j \neq i$ .

Если слова нет в словаре (OOV – out of vocabulary), то для него зарезервирована отдельная позиция в данном векторе:  $x_{|V|+1} = 1$

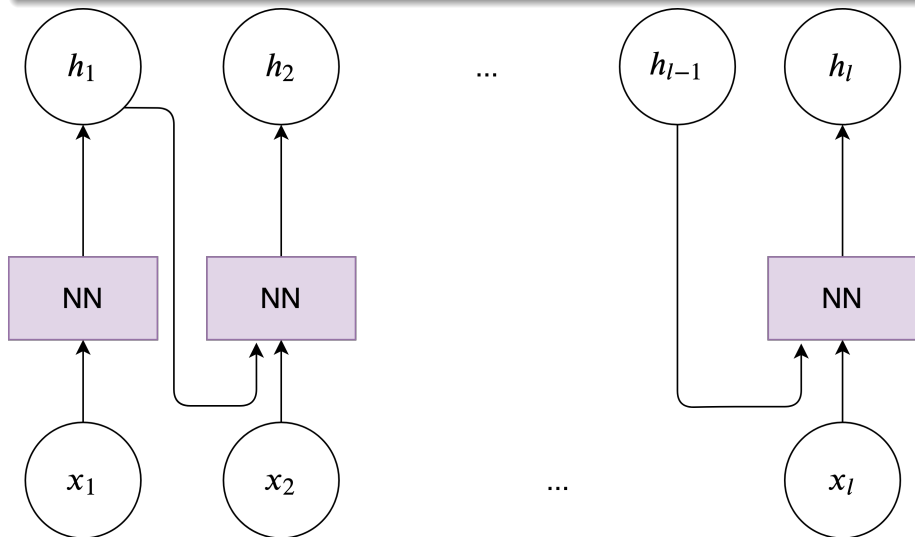
Тогда текст можно представить как последовательность векторов словоформ  $x_1, x_2, \dots, x_l$ ,  $x_i \in \mathbb{R}^d$ .

Как получить представление каждого слова, с учетом его нынешнего контекста, чтобы его классифицировать на предмет вхождения в сущность?

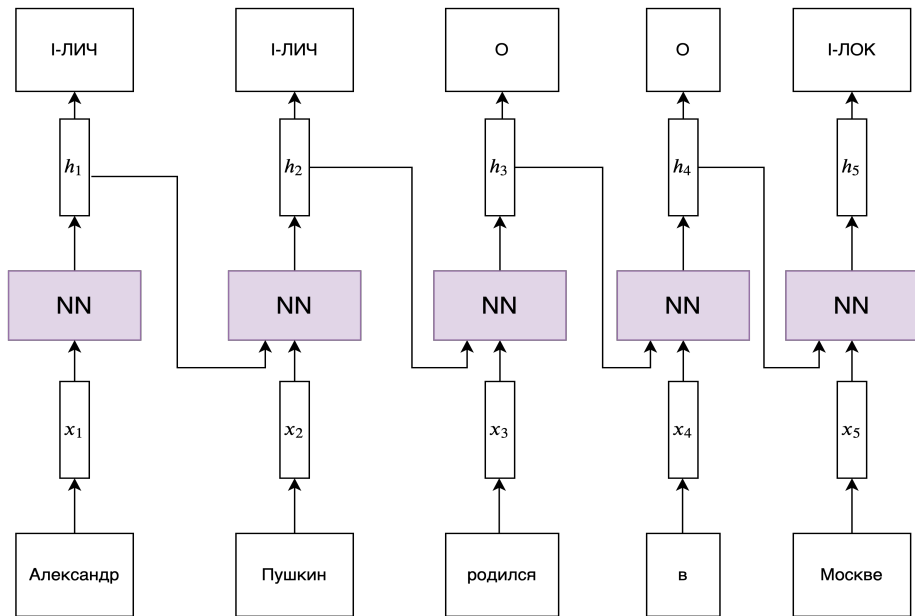
- сверточные нейронные сети (CNN)
- нейронные сети на основе механизма внимания (Transformer)
- рекуррентные нейронные сети (GRU, LSTM)

# Рекуррентные нейронные сети

Формально:  $\mathbf{h}_t = f(\mathbf{x}_t, \mathbf{h}_{t-1})$ ,  $\mathbf{h}_0 = \mathbf{0}$ ,  $t \in [1, l]$



# Рекуррентные нейронные сети



# Простейшая рекуррентная нейронная сеть

- $h_0 = 0$
- $h_t = f(W_x x_t + W_h h_{t-1} + b)$



# Проблема catastrophic forgetting

- Пусть на вход дано 3 вектора  $x_1, x_2, x_3$
- Развернем выход сети при  $b = 0$ :

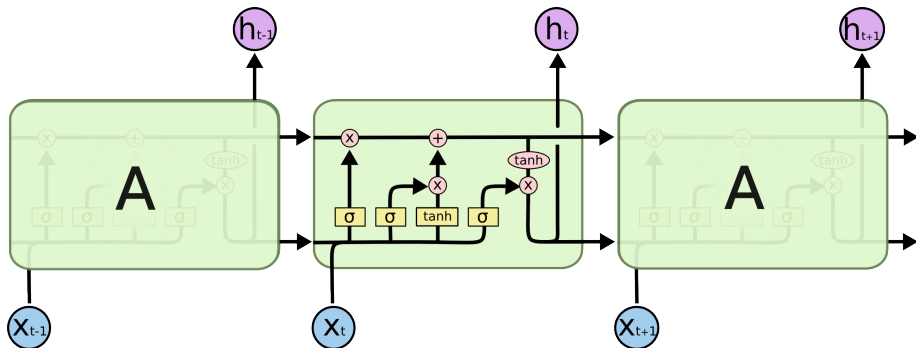
$$h_3 = f(W_x x_3 + W_h f(W_x x_2 + W_h h_1))$$

- Считая, что мы работаем со скалярами:

$$\frac{\partial h_3}{\partial h_1} = f'(W_x x_3 + W_h f(W_x x_2 + W_h h_1)) * W_h * f'(W_x x_2 + W_h h_1) * W_h$$

- Предыдущий выход не влияет на текущее решение нашей сети, а значит наша сеть забывает, что было ранее :(

# LSTM



<sup>0</sup><https://colah.github.io/posts/2015-08-Understanding-LSTMs/>

- $\mathbf{f}_t = \sigma(\mathbf{W}_f \mathbf{x}_t + \mathbf{U}_f \mathbf{h}_{t-1} + \mathbf{b}_f)$  – что забываем  
 $\mathbf{i}_t = \sigma(\mathbf{W}_i \mathbf{x}_t + \mathbf{U}_i \mathbf{h}_{t-1} + \mathbf{b}_i)$  – что запоминаем  
 $\mathbf{o}_t = \sigma(\mathbf{W}_o \mathbf{x}_t + \mathbf{U}_o \mathbf{h}_{t-1} + \mathbf{b}_o)$  – что выдаем  
 $\tilde{\mathbf{c}}_t = \tanh(\mathbf{W}_c \mathbf{x}_t + \mathbf{U}_c \mathbf{h}_{t-1} + \mathbf{b}_c)$   
 $\mathbf{c}_t = \mathbf{f}_t \odot \mathbf{c}_{t-1} + \mathbf{i}_t \odot \tilde{\mathbf{c}}_t$   
 $\mathbf{h}_t = \mathbf{o}_t \odot \tanh(\mathbf{c}_t)$
- $\frac{\partial \mathbf{c}_t}{\partial \mathbf{c}_{t-1}} = \text{diag}(\mathbf{f}_t)$  – когда сеть не хочет забывать, градиент не затухает

В задачах обработки текстов важен не только левый контекст словоформы, но и правый:

**Александр Пушкин родился** в Москве

Поэтому используют BiLSTM – две параллельные LSTM, первой на вход подается исходная последовательность, второй – исходная последовательность в обратном порядке. Их выходы  $\vec{h}_t$  и  $\overleftarrow{h}_t$  конкатенируются.

Мы закодировали исходную последовательность векторов словоформ  $x_1, \dots, x_l$  в последовательность векторов  $h_1, \dots, h_l$ , которые учитывают контекст.

Если мы хотим классифицировать текст, то нужно получить фиксированный вектор  $a$ , который далее отправится в классификатор:

- взять последний выход  $a = h_l$
- mean pooling  $a_i = \frac{\sum_j h_{j,i}}{l}$
- max pooling  $a_i = \max_j h_{j,i}$
- механизм внимания

# Что посмотреть, если интересно

- A.Karpathy CS231n Winter 2016 lectures
- C.Manning, R.Socher CS224n Winter 2017 lectures
- Воронцов К.В. Машинное обучение

Вопросы?