

Когда нейросеть устойчива: оценки Липшицевости

Сюжет: когда панда становится гibbonом

В 2014 г. Иэн Гудфеллоу и соавторы опубликовали ставшую знаменитой картинку (рис. 0.10): обычная фотография панды, которую современная нейросеть уверенно опознаёт как «panda» с вероятностью 57,7%. К каждому пикселю прибавлена крошечная, визуально незаметная человеку поправка. На обновлённой картинке сеть теперь уверенно (с вероятностью 99,3%) считает, что видит *гibbonа*.

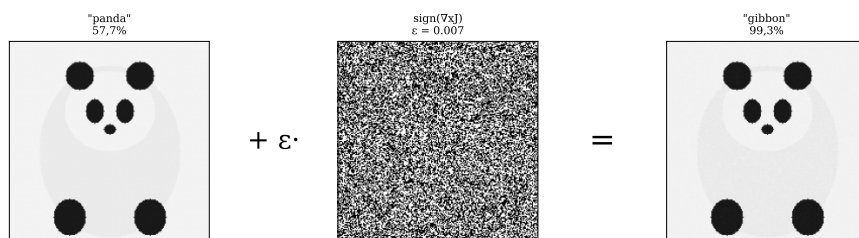


Рис. 0.10. Adversarial-атака на нейросеть (Goodfellow et al., 2014). Слева — оригинал, классифицирован как «panda» с уверенностью 57,7%. В центре — почти невидимое возмущение, масштабированное для наглядности. Справа — атакованный пример, классифицируется как «gibbon» с уверенностью 99,3%.

Этот эффект называется **adversarial-атакой** и долгое время считался экзотикой. Сейчас понятно: причина в простом числе, которое сопровождает любую функцию — её *константе Липшица*.

Константа Липшица: формальное определение

i Определение 0.5. Константа Липшица функции

Функция $f : \mathbb{R}^n \rightarrow \mathbb{R}^m$ называется **Липшицевой** с константой $L > 0$, если для любых $x, y \in \mathbb{R}^n$ выполнено

$$\|f(x) - f(y)\| \leq L \|x - y\|. \quad (1)$$

Наименьшее L , для которого (0.9) выполнено, называется *константой Липшица* функции f и обозначается $\text{Lip}(f)$.

Содержательно $\text{Lip}(f)$ показывает, во сколько раз может *усилиться* входное возмущение на выходе функции. Если

$\text{Lip}(f) = L$, то малое возмущение δ на входе может вызвать возмущение до $L \cdot \|\delta\|$ на выходе.

Пример. Для линейного отображения $f(x) = Wx$ имеем $f(x) - f(y) = W(x - y)$, откуда $\|f(x) - f(y)\| \leq \|W\|_2 \cdot \|x - y\|$, и константа Липшица — это в точности спектральная норма матрицы W :

$$\text{Lip}(Wx) = \|W\|_2 = \sigma_1(W). \quad (2)$$

Здесь $\sigma_1(W)$ — наибольшее сингулярное число матрицы W . Вот *здесь* в нашу историю входит SVD из § 0.3.

Композиция: Липшицева константа цепочки

⚠ Теорема 0.6. Композиция Липшицевых функций

Если f и g — Липшицевы с константами L_f и L_g , то их композиция $f \circ g$ — тоже Липшицева, причём

$$\text{Lip}(f \circ g) \leq \text{Lip}(f) \cdot \text{Lip}(g). \quad (3)$$

Доказательство. $\|f(g(x)) - f(g(y))\| \leq L_f \|g(x) - g(y)\| \leq L_f L_g \|x - y\|$. $\square$

Применение к нейросети

Возьмём типичную нейросеть глубины K :

$$f(x) = \sigma(W_K \sigma(W_{K-1} \sigma(\dots \sigma(W_1 x))), \quad (4)$$

где W_i — матрицы весов, σ — поэлементная активация (ReLU, сигмоида, и т.п.). Для большинства активаций $\text{Lip}(\sigma) \leq 1$ (для ReLU и tanh она равна 1, а для логистической сигмоиды — 1/4). По теореме 0.6:

$$\text{Lip}(f) \leq \prod_{i=1}^K \|W_i\|_2 \quad (5)$$

Это и есть тот «лишний рычаг», который объясняет уязвимость нейросетей. Если $\|W_i\|_2 \approx 2$ для каждого слоя из $K = 20$ слоёв, то возмущение на входе *потенциально* усиливается в $2^{20} \approx 10^6$ раз. Невидимая поправка δ с $\|\delta\| = 10^{-3}$ может

вызвать изменение выхода до 10^3 — этого хватит, чтобы перебросить классификатор из одного класса в другой.

Как защищаться: спектральная нормализация

Естественная стратегия: после каждого шага обучения нормализовать матрицу W_i так, чтобы $\|W_i\|_2 = 1$. Это и есть метод **Spectral Normalization** (Миято и соавт., 2018), сделавший обучение генеративных нейросетей (GAN) намного стабильнее.

Технически: на каждом шаге считаем $\sigma_1(W)$ (с помощью одной итерации степенного метода — почти бесплатно!) и заменяем $W \mapsto W/\sigma_1(W)$.

```
def spectral_normalize(W, n_iter=1):
    """One step of power iteration to estimate sigma_1(W) and rescale."""
    u = np.random.randn(W.shape[0])
    for _ in range(n_iter):
        v = W.T @ u; v /= np.linalg.norm(v)
        u = W @ v; u /= np.linalg.norm(u)
    sigma_1 = u @ W @ v
    return W / sigma_1
```

💡 Это интересно

Связь с устойчивостью в физике. Похожая ситуация хорошо известна в теории динамических систем: если константа Липшица отображения > 1 , малое возмущение начального условия экспоненциально нарастает (классический эффект бабочки в погодных моделях). Spectral Normalization можно рассматривать как искусственное гашение этого эффекта.