



مرکز آموزش‌های الکترونیکی دانشگاه شهید بهشتی

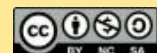
شبکه‌های عصبی مصنوعی

Artificial Neural Networks

Ali Katanforoush*

* Department of Data and Computer Sciences, Shahid Beheshti University

1 / 23



Weight Regularization منظم سازی وزن

ضرب منظم سازی
L2-norm

$$\text{Loss}(w) = \frac{1}{N} \sum_{i=1}^N (y^{(i)} - \hat{y}(x^{(i)}; w))^2 + \frac{\lambda}{2} \sum_{j=1}^M w_j^2$$

اضافه کردن یک عبارت منظم سازی در تعریف تابع زیان

Neural networks

Training neural networks - regularization

$$w^{\text{new}} = (1 - \alpha) w^{\text{old}} - \eta \nabla_w f$$

weight decay

REGULARIZATION

Topics: L2 regularization

$$\Omega(\boldsymbol{\theta}) = \sum_k \sum_i \sum_j \left(w_{i,j}^{(k)} \right)^2 = \sum_k \| \mathbf{w}^{(k)} \|_F^2$$

- Gradient: $\nabla_{\mathbf{w}^{(k)}} \Omega(\boldsymbol{\theta}) = 2 \mathbf{w}^{(k)}$
- Only applied on weights, not on biases (weight decay)
- Can be interpreted as having a Gaussian prior over the weights

REGULARIZATION

Topics: L1 regularization

$$\Omega(\theta) = \sum_k \sum_i \sum_j |w_{i,j}^{(k)}|$$

- Gradient: $\nabla_{\mathbf{w}^{(k)}} \Omega(\theta) = \text{sign}(\mathbf{w}^{(k)})$
 - where $\text{sign}(\mathbf{w}^{(k)})_{i,j} = 1_{\mathbf{w}_{i,j}^{(k)} > 0} - 1_{\mathbf{w}_{i,j}^{(k)} < 0}$
- Also only applied on weights
- Unlike L2, L1 will push certain weights to be exactly 0
- Can be interpreted as having a Laplacian prior over the weights

Neural networks

Deep learning - dropout

DEEP LEARNING

Topics: why training is hard

- Depending on the problem, one or the other situation will tend to prevail

کم برازی

- If first hypothesis (underfitting): use better optimization

- this is an active area of research

بہتر برازی

- If second hypothesis (overfitting): use better regularization

- unsupervised learning

- stochastic «dropout» training

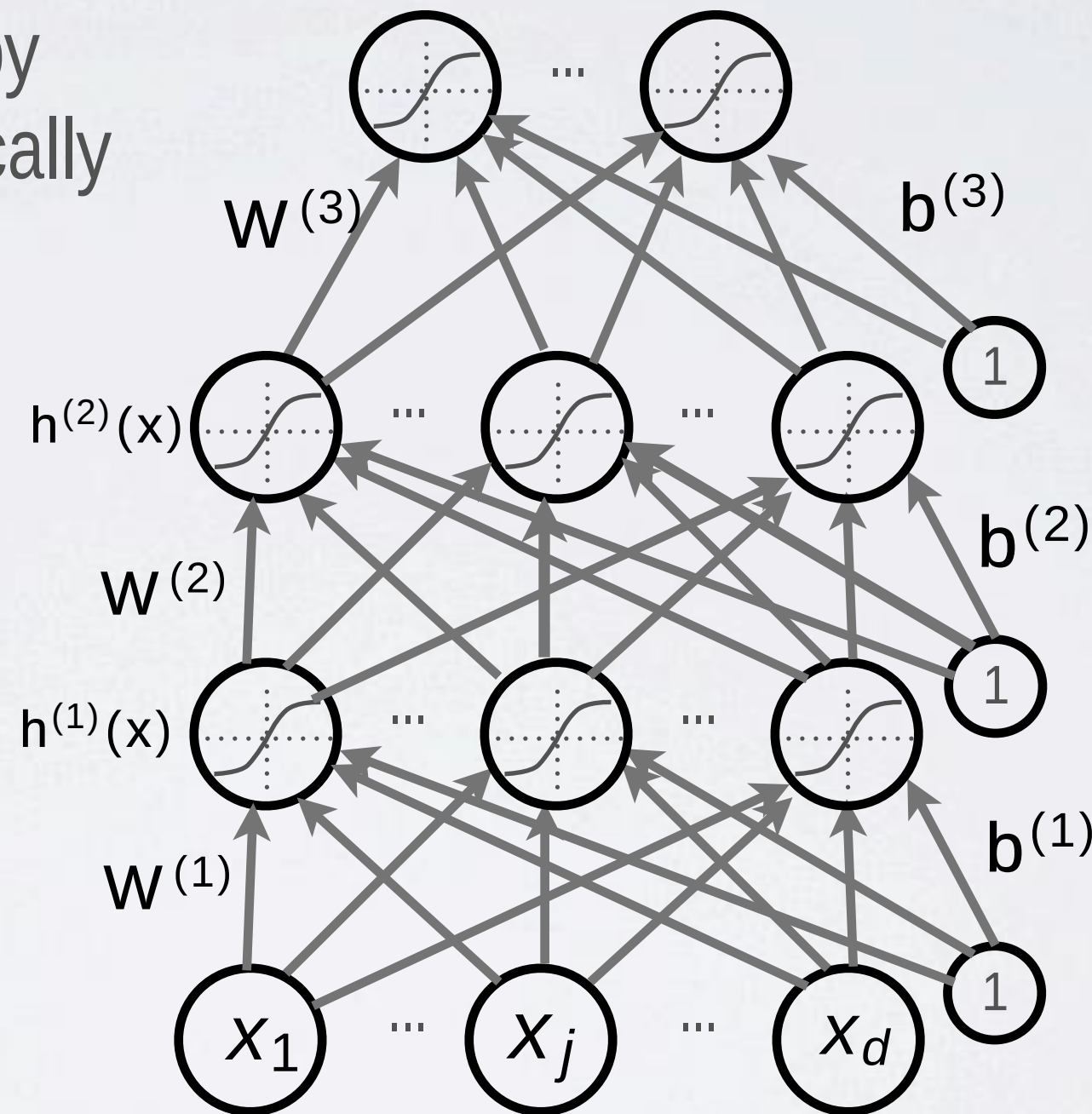
ریز

→ (GAN)

DROPOUT

Topics: dropout

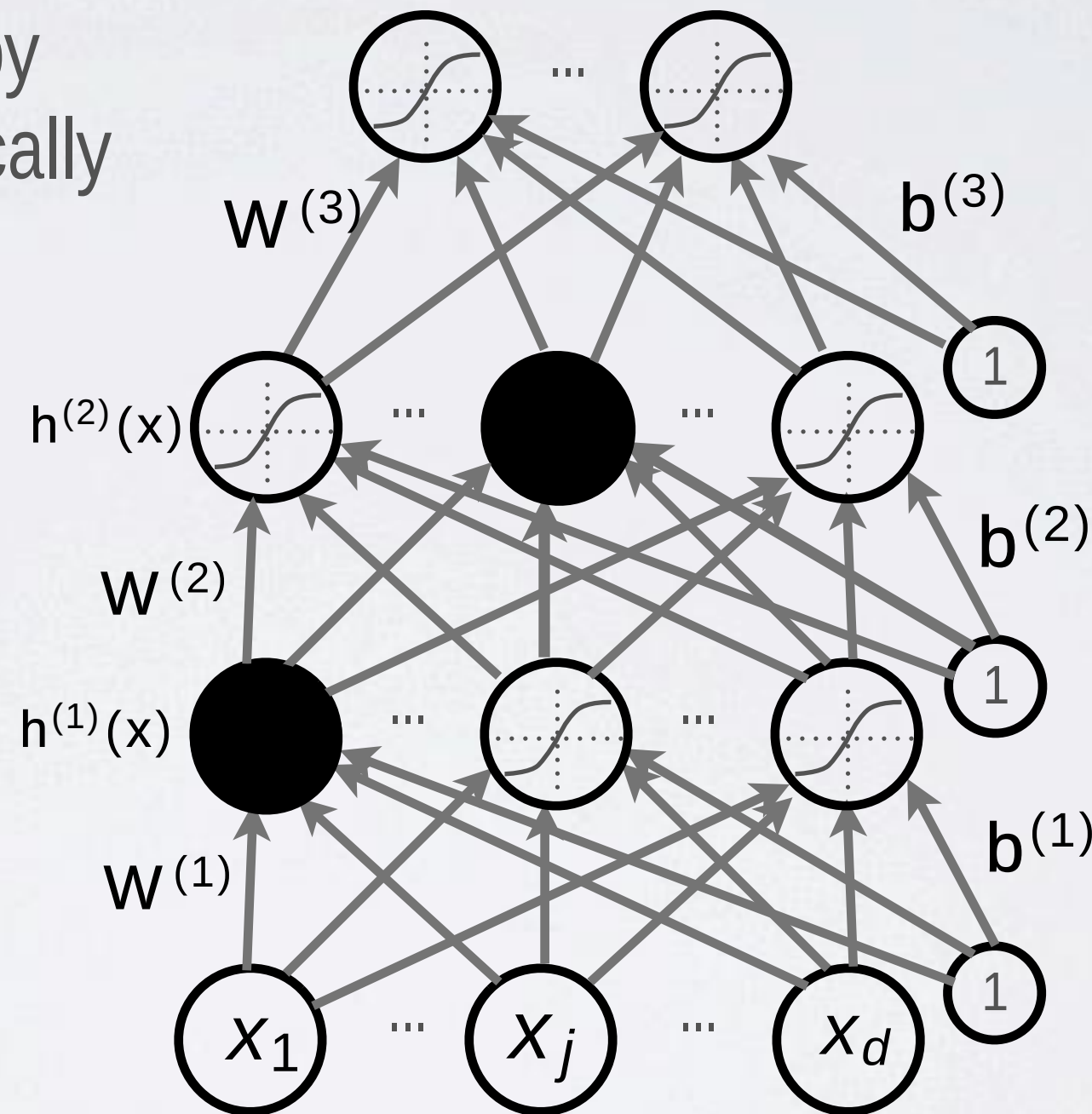
- Idea: «cripple» neural network by removing hidden units stochastically
 - each hidden unit is set to 0 with probability 0.5
 - hidden units cannot co-adapt to other units
 - hidden units must be more generally useful
- Could use a different dropout probability, but 0.5 usually works well



DROPOUT

Topics: dropout

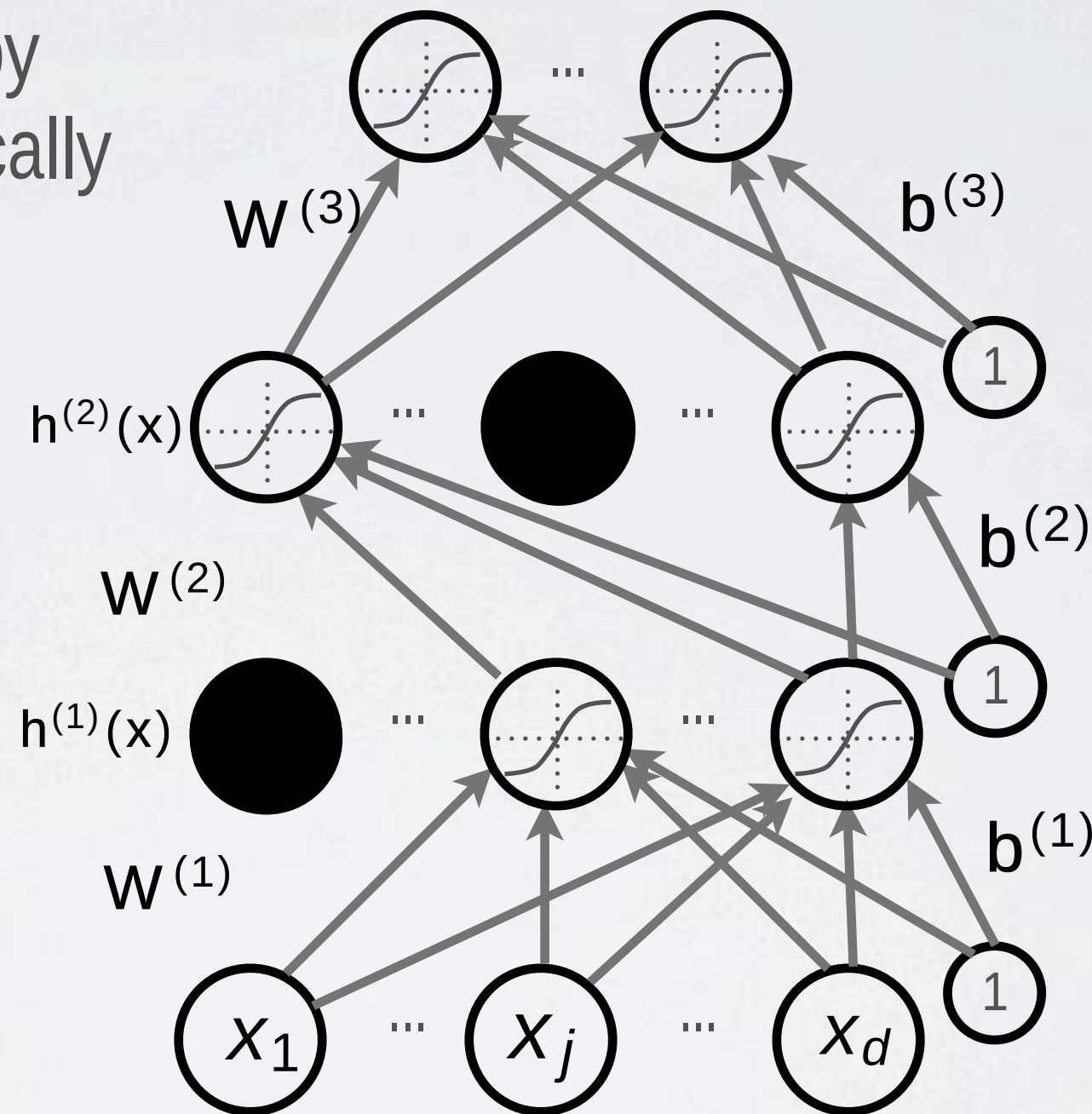
- Idea: «cripple» neural network by removing hidden units stochastically
 - each hidden unit is set to 0 with probability 0.5
 - hidden units cannot co-adapt to other units
 - hidden units must be more generally useful
- Could use a different dropout probability, but 0.5 usually works well



DROPOUT

Topics: dropout

- Idea: «cripple» neural network by removing hidden units stochastically
 - each hidden unit is set to 0 with probability 0.5
 - hidden units cannot co-adapt to other units
 - hidden units must be more generally useful
- Could use a different dropout probability, but 0.5 usually works well



DROPOUT

Topics: dropout

- Use random binary masks $\mathbf{m}^{(k)}$

- layer pre-activation for $k > 0$ ($h^{(0)}(x) = x$)

$$\mathbf{a}^{(k)}(x) = \mathbf{b}^{(k)} + \mathbf{W}^{(k)} \mathbf{h}^{(k-1)}(x)$$

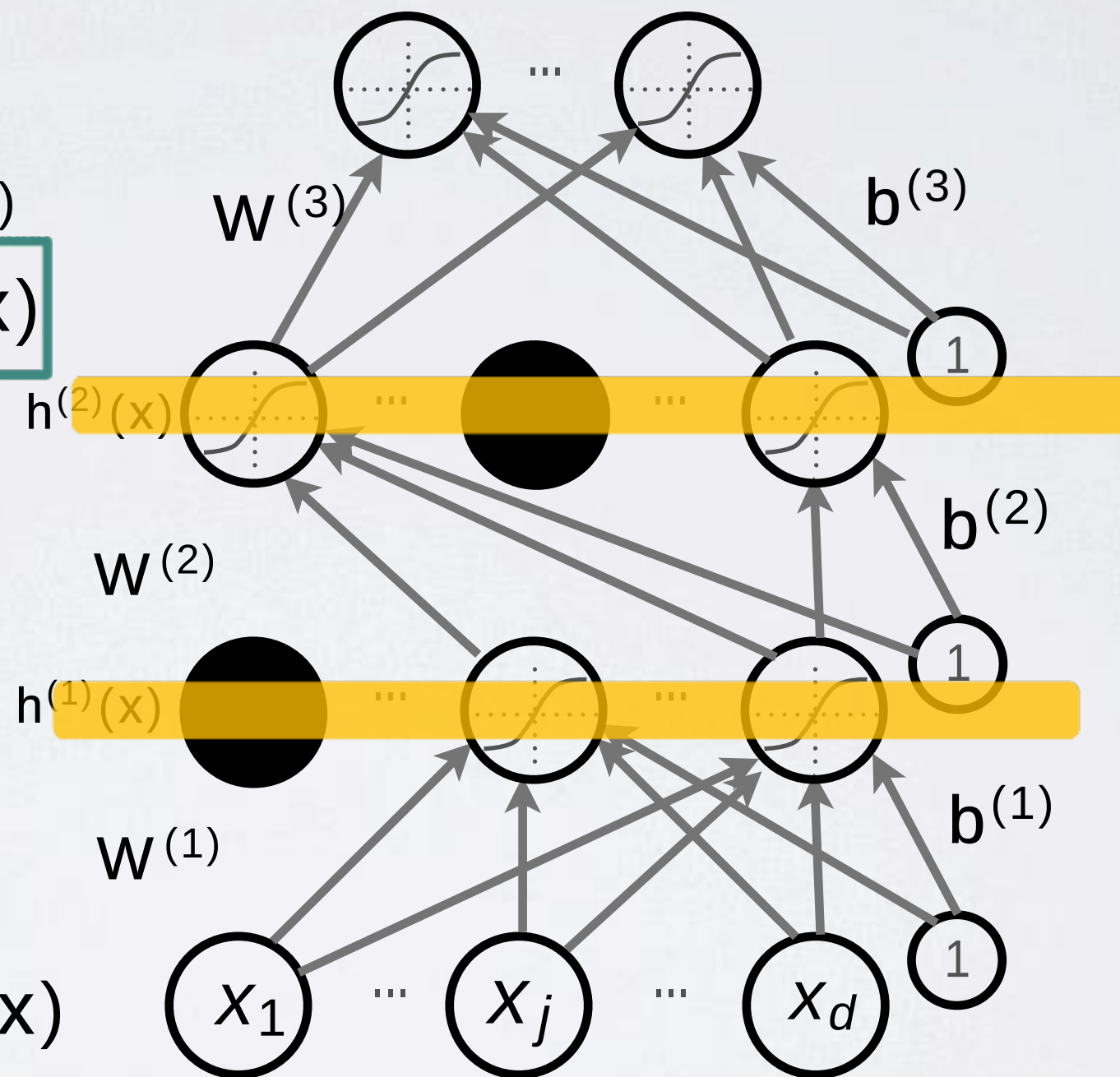
- hidden layer activation (k from 1 to L):

$$\mathbf{h}^{(k)}(x) = \mathbf{g}(\mathbf{a}^{(k)}(x))$$

بدون لایه Dropout

- output layer activation ($k = L + 1$):

$$\mathbf{h}^{(L+1)}(x) = \mathbf{o}(\mathbf{a}^{(L+1)}(x)) = \mathbf{f}(x)$$



DROPOUT

$$m = \langle \underline{0}, \underline{1}, \underline{1}, \underline{0} \dots, \underline{1} \rangle$$

Topics: dropout

- Use random binary masks $m^{(k)}$

- layer pre-activation for $k > 0$ ($h^{(0)}(x) = x$)

$$a^{(k)}(x) = b^{(k)} + W^{(k)} h^{(k-1)}(x)$$

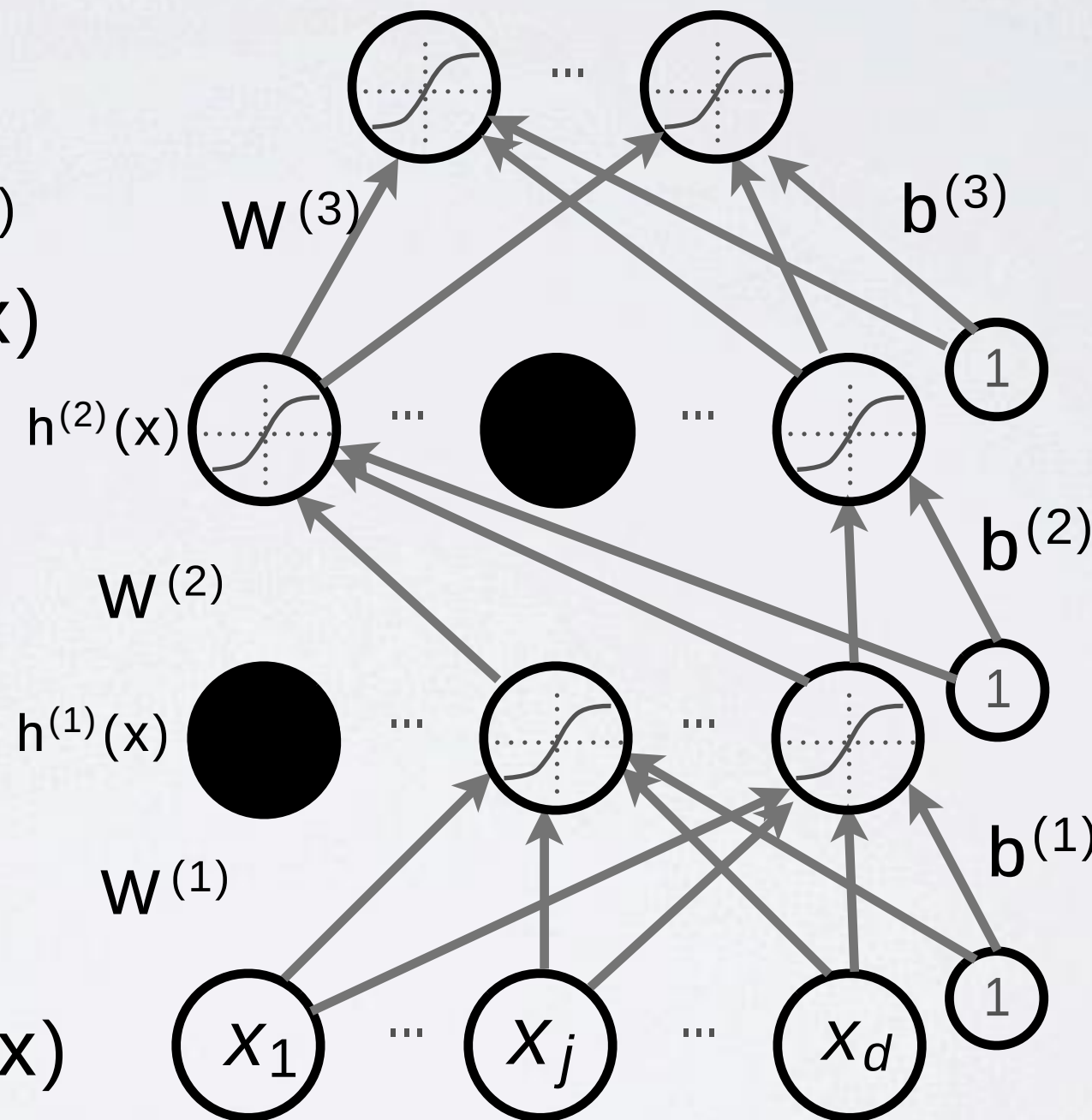
Dropout ≈ 1 !

- hidden layer activation (k from 1 to L):

$$h^{(k)}(x) = g(a^{(k)}(x)) \odot m^{(k)}$$

- output layer activation ($k = L + 1$):

$$h^{(L+1)}(x) = o(a^{(L+1)}(x)) = f(x)$$



$$0 \leq 1$$

dropout ϵ_j

$$p = 0.2$$

$$E[m_j] =$$

$$0 \times p + 1 \times (1 - p)$$

$$= 1 - p = 0.8$$

DROPOUT

Topics: dropout backpropagation

- This assumes a forward propagation has been made before

- compute output gradient (before activation)

$$\nabla_{\mathbf{a}^{(L+1)}(\mathbf{x})} - \log f(\mathbf{x})_y \Leftarrow -(\mathbf{e}(y) - \mathbf{f}(\mathbf{x}))$$

- for k from $L+1$ to 1

- compute gradients of hidden layer parameter

$$\nabla_{\mathbf{W}^{(k)}} - \log f(\mathbf{x})_y \Leftarrow (\nabla_{\mathbf{a}^{(k)}(\mathbf{x})} - \log f(\mathbf{x})_y) \mathbf{h}^{(k-1)}(\mathbf{x})^\top$$

$$\nabla_{\mathbf{b}^{(k)}} - \log f(\mathbf{x})_y \Leftarrow \nabla_{\mathbf{a}^{(k)}(\mathbf{x})} - \log f(\mathbf{x})_y$$

- compute gradient of hidden layer below

$$\nabla_{\mathbf{h}^{(k-1)}(\mathbf{x})} - \log f(\mathbf{x})_y \Leftarrow \mathbf{W}^{(k)\top} (\nabla_{\mathbf{a}^{(k)}(\mathbf{x})} - \log f(\mathbf{x})_y)$$

- compute gradient of hidden layer below (before activation)

$$\nabla_{\mathbf{a}^{(k-1)}(\mathbf{x})} - \log f(\mathbf{x})_y \Leftarrow (\nabla_{\mathbf{h}^{(k-1)}(\mathbf{x})} - \log f(\mathbf{x})_y) \odot [\dots, g'(a^{(k-1)}(\mathbf{x})_j), \dots]$$

بدون لايه Dropout

DROPOUT

Topics: dropout backpropagation

- This assumes a forward propagation has been made before

- compute output gradient (before activation)

$$\nabla_{\mathbf{a}^{(L+1)}(\mathbf{x})} - \log f(\mathbf{x})_y \Leftarrow -(\mathbf{e}(y) - \mathbf{f}(\mathbf{x}))$$

- for k from $L+1$ to 1

- compute gradients of hidden layer parameter

$$\nabla_{\mathbf{W}^{(k)}} - \log f(\mathbf{x})_y \Leftarrow (\nabla_{\mathbf{a}^{(k)}(\mathbf{x})} - \log f(\mathbf{x})_y) \mathbf{h}^{(k-1)}(\mathbf{x})^\top$$

$$\nabla_{\mathbf{b}^{(k)}} - \log f(\mathbf{x})_y \Leftarrow \nabla_{\mathbf{a}^{(k)}(\mathbf{x})} - \log f(\mathbf{x})_y$$

- compute gradient of hidden layer below

$$\nabla_{\mathbf{h}^{(k-1)}(\mathbf{x})} - \log f(\mathbf{x})_y \Leftarrow \mathbf{W}^{(k)\top} (\nabla_{\mathbf{a}^{(k)}(\mathbf{x})} - \log f(\mathbf{x})_y)$$

- compute gradient of hidden layer below (before activation)

$$\nabla_{\mathbf{a}^{(k-1)}(\mathbf{x})} - \log f(\mathbf{x})_y \Leftarrow (\nabla_{\mathbf{h}^{(k-1)}(\mathbf{x})} - \log f(\mathbf{x})_y) \odot [\dots, g'(a^{(k-1)}(\mathbf{x})_j), \dots] \odot \mathbf{m}^{(k-1)}$$

includes the
mask $\mathbf{m}^{(k-1)}$

Dropout !!!

DROPOUT

Topics: test time classification

- At test time, we replace the masks by their expectation
 - this is simply the constant vector 0.5 if dropout probability is 0.5
 - for single hidden layer, can show this is equivalent to taking the geometric average of all neural networks, with all possible binary masks
- Can be combined with unsupervised pre-training
- Beats regular backpropagation on many datasets
 - Improving neural networks by preventing co-adaptation of feature detectors.
Hinton, Srivastava, Krizhevsky, Sutskever and Salakhutdinov, 2012.

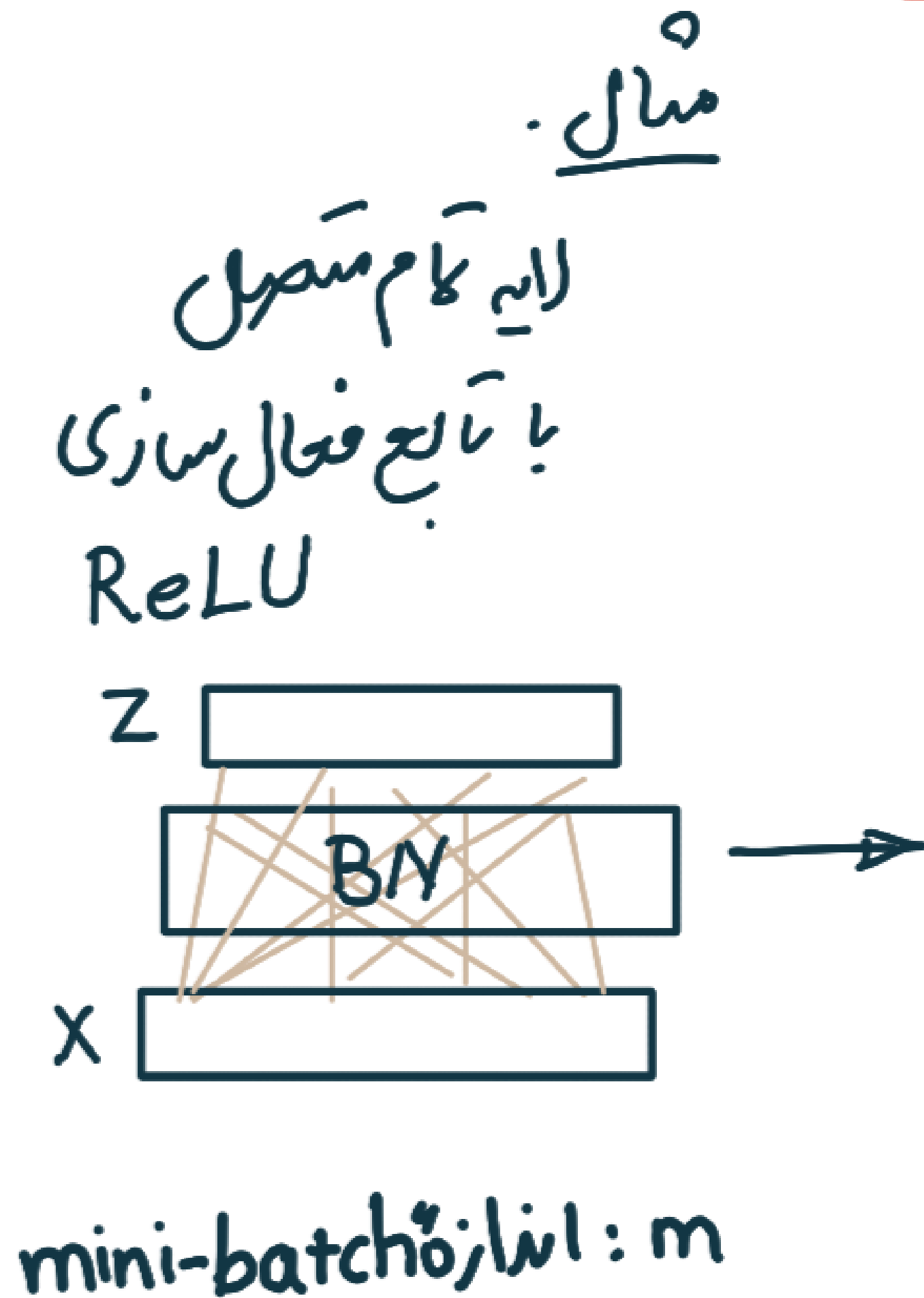
Batch Normalization

$$Z = XW$$

$$\tilde{Z} = Z - \frac{1}{m} \sum_{i=1}^m Z_{i,:}$$

$$\hat{Z} = \frac{\tilde{Z}}{\sqrt{\epsilon + \frac{1}{m} \sum_{i=1}^m \tilde{Z}_{i,:}^2}}$$

$$H = \max\{0, \gamma \hat{Z} + \beta\}$$



“Batch Normalization: Accelerating Deep Network Training by Reducing Internal Covariate Shift,” Ioffe and Szegedy 2015