



CTU

CZECH TECHNICAL
UNIVERSITY
IN PRAGUE

Deep Learning Essentials

6. Activation, Normalization and Regularization

Activation Functions, BatchNorm, Dropout, weight decay

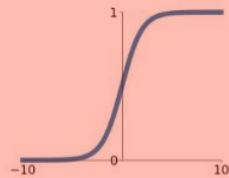
Lukáš Neumann

Adapted from [B3B33UROB](#) slides of Karel Zimmerman

Activation Functions

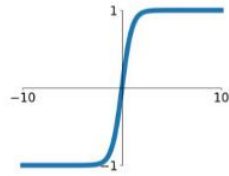
Sigmoid

$$\sigma(x) = \frac{1}{1+e^{-x}}$$



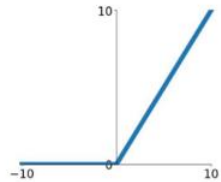
tanh

$$\tanh(x)$$



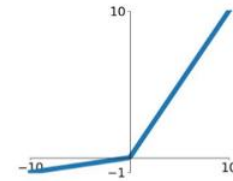
ReLU

$$\max(0, x)$$



Leaky ReLU

$$\max(0.1x, x)$$

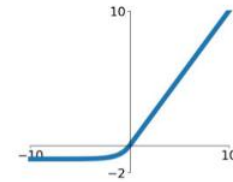


Maxout

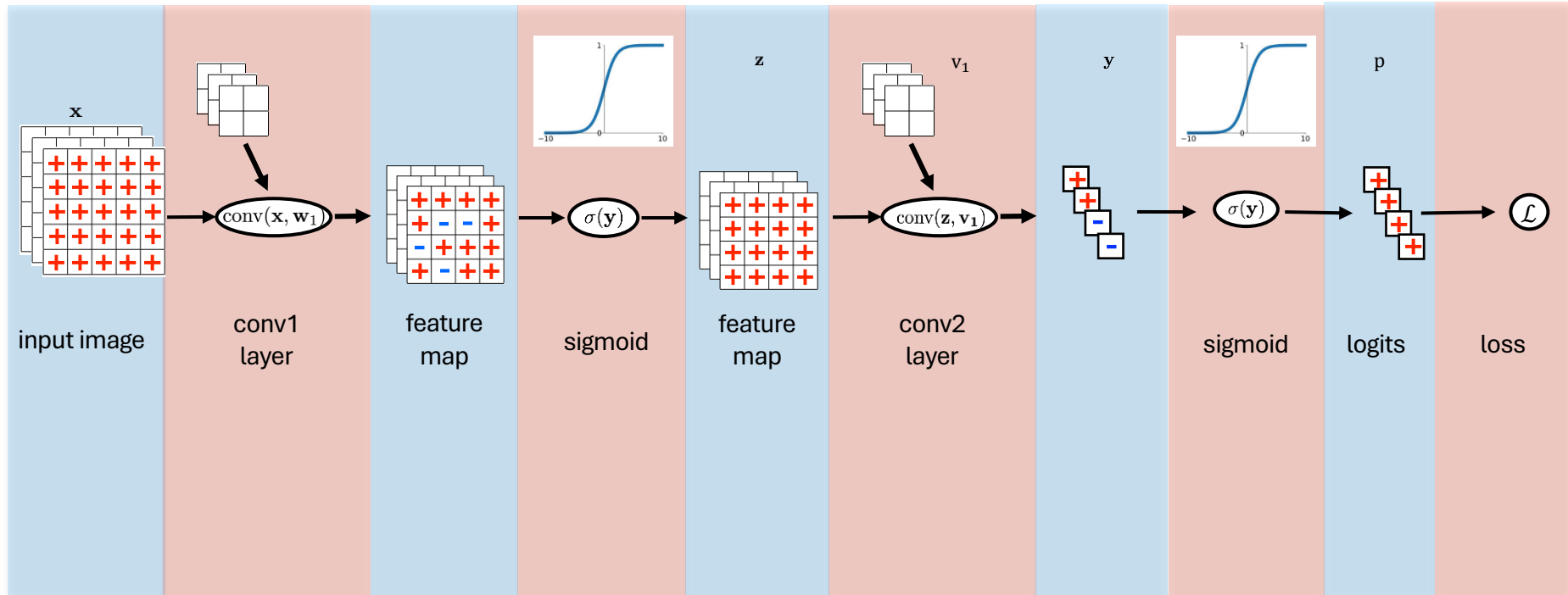
$$\max(w_1^T x + b_1, w_2^T x + b_2)$$

ELU

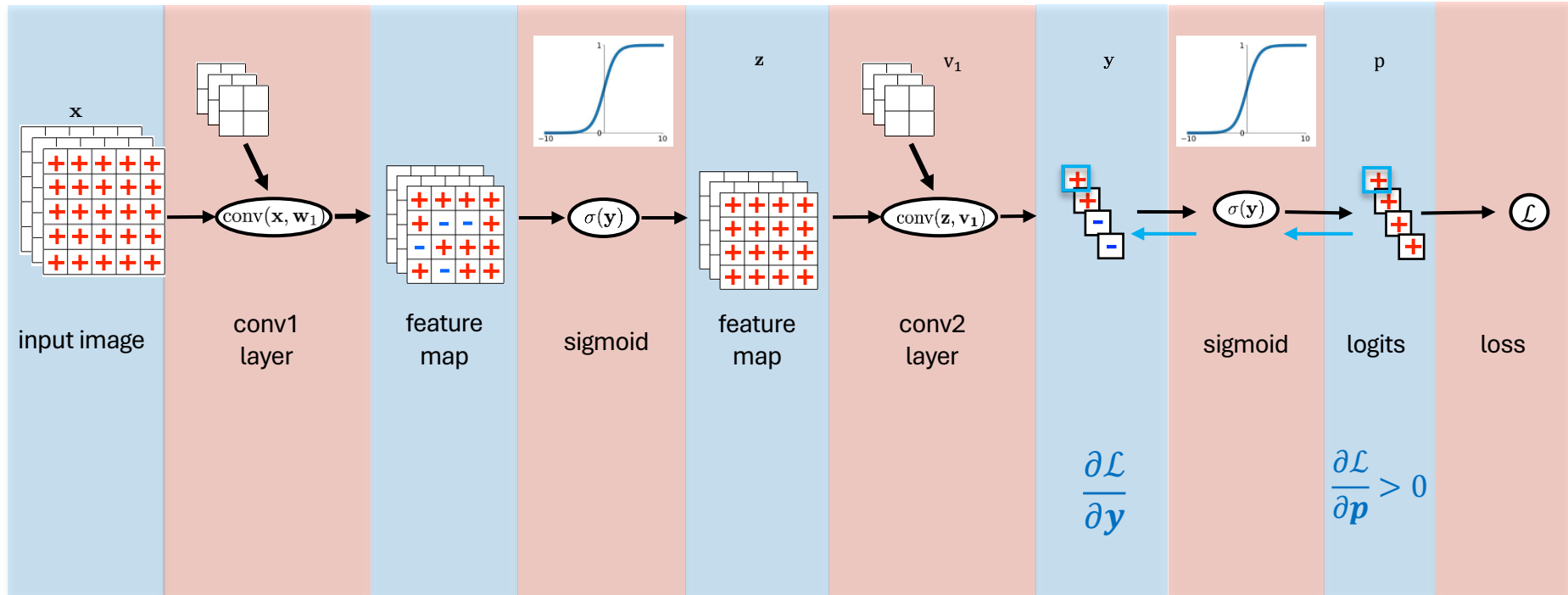
$$\begin{cases} x & x \geq 0 \\ \alpha(e^x - 1) & x < 0 \end{cases}$$



Sigmoid



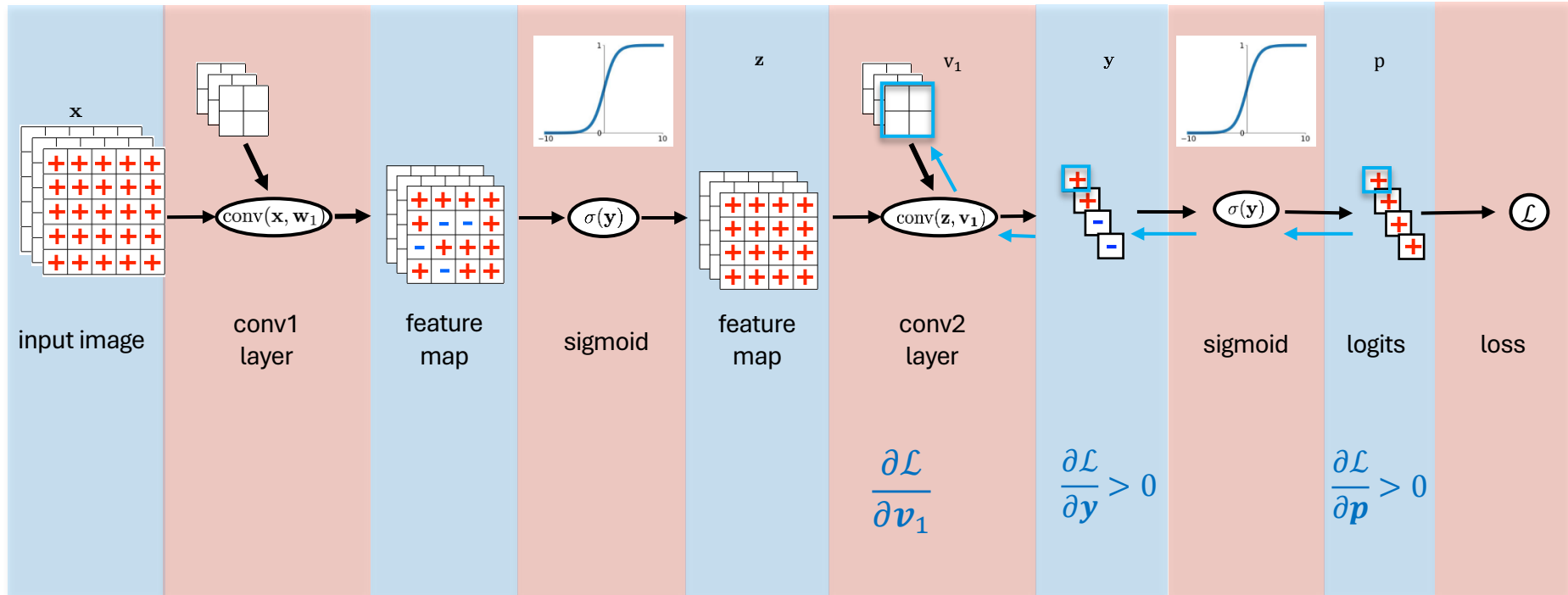
Sigmoid



$$\boxed{\frac{\partial \mathcal{L}}{\partial y}} = \frac{\partial p}{\partial y} \cdot \frac{\partial \mathcal{L}}{\partial p} = \boxed{\frac{\partial \sigma(y)}{\partial y}} \cdot \boxed{\frac{\partial \mathcal{L}}{\partial p}}$$

> 0 > 0 > 0

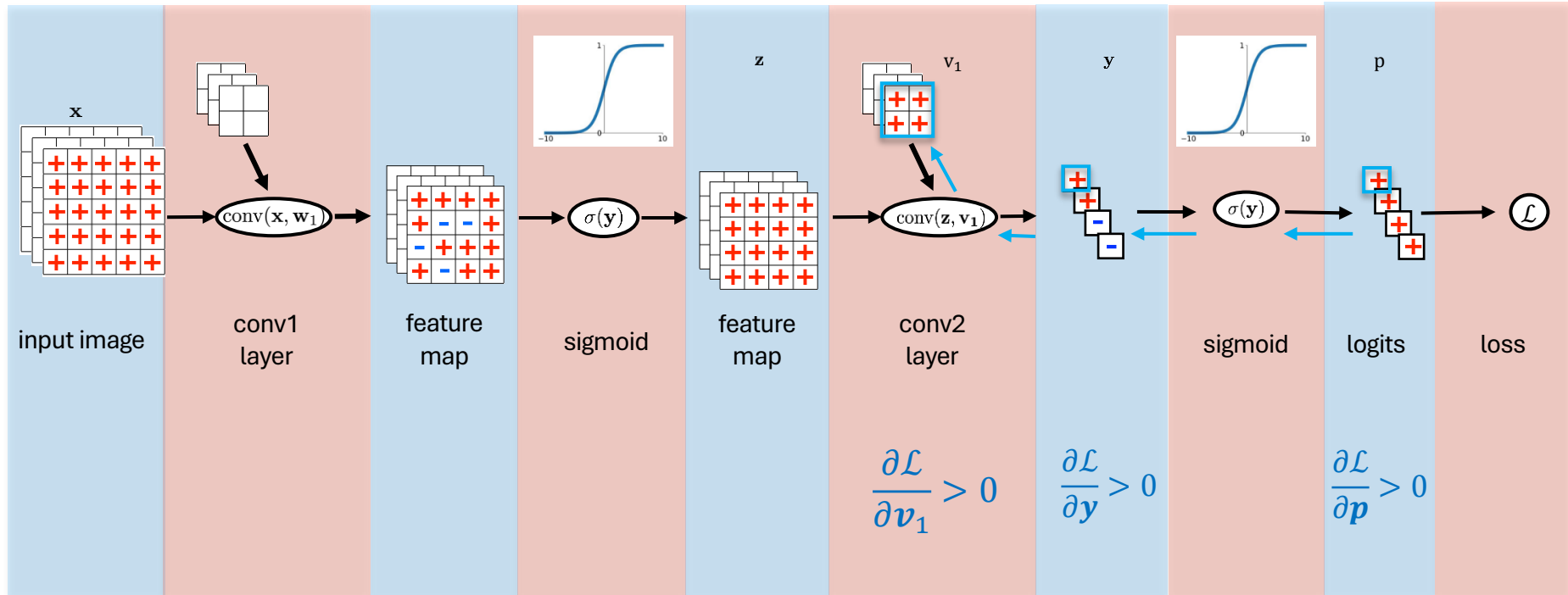
Sigmoid



$$\frac{\partial \mathcal{L}}{\partial v_1} \approx \text{conv}\left(\frac{\partial \mathcal{L}}{\partial y}, z\right)$$

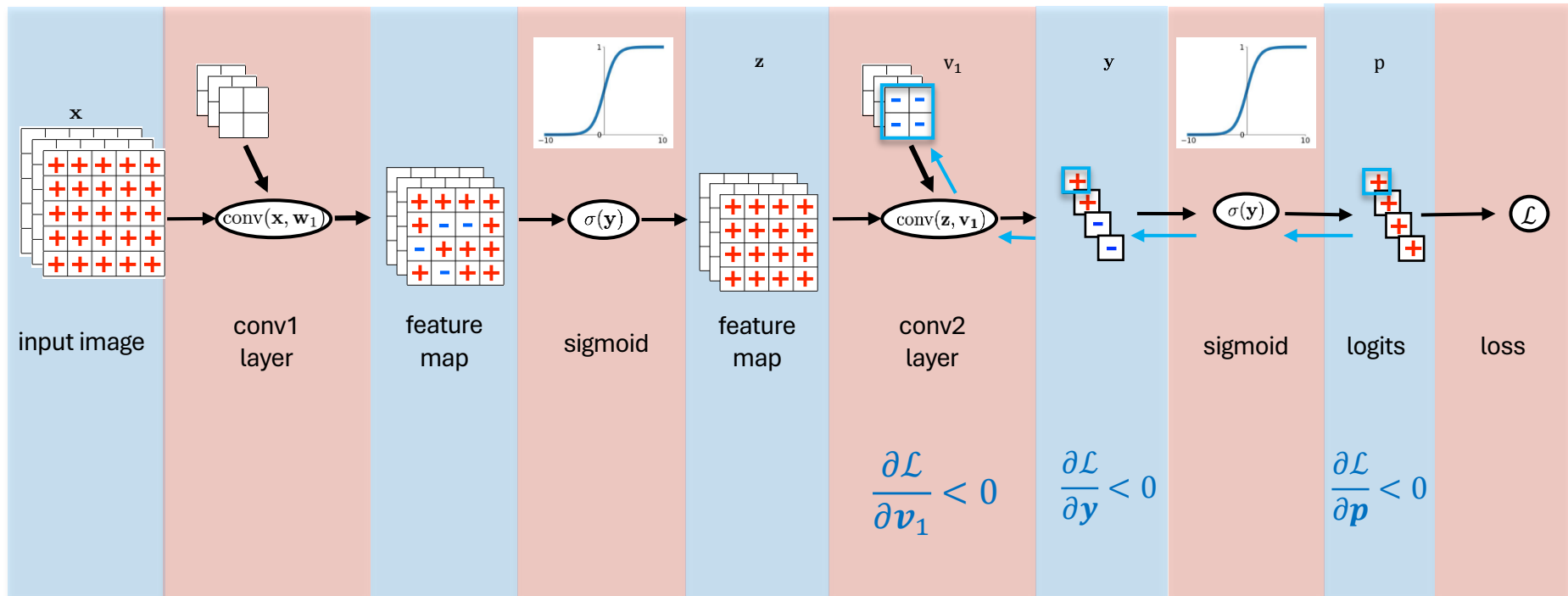
> 0 > 0 > 0

Sigmoid

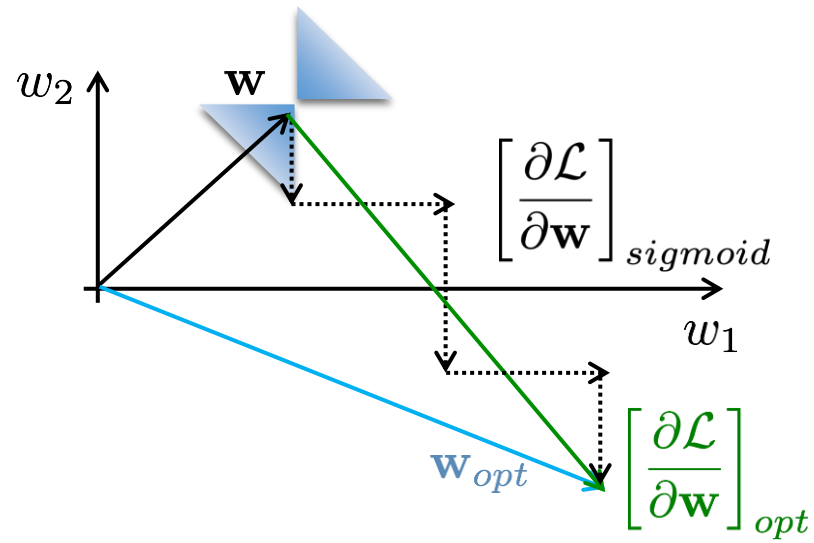


Sigmoid

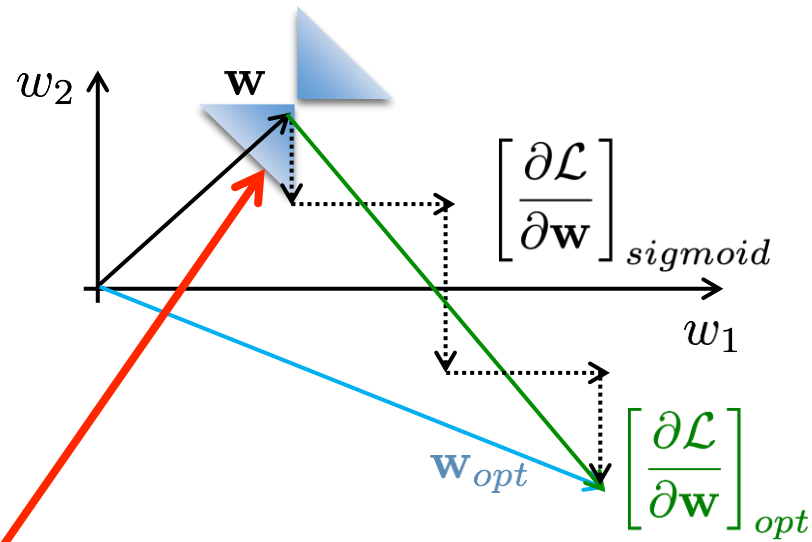
- Upstream gradient is positive $\rightarrow$ weights gradient is positive for **all weights**
- Upstream gradient is negative $\rightarrow$ weights gradient is negative for **all weights**



Sigmoid



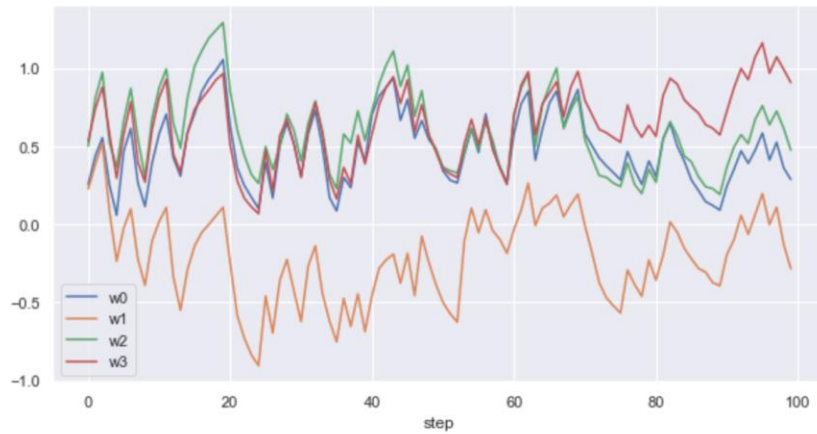
Sigmoid



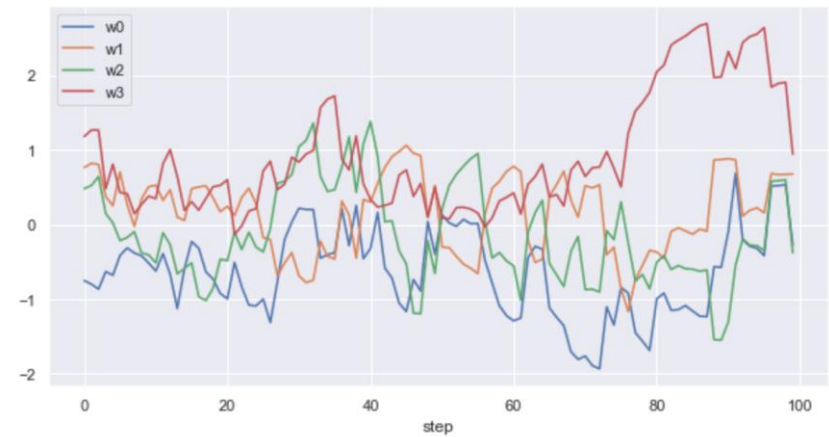
How big fraction is the blue region in a d -dimensional space? $\frac{2}{2^d}$

Sigmoid

Sigmoid Activation



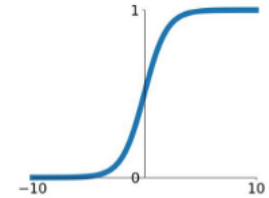
Tanh Activation



Sigmoid

- No gradient when saturated
- Not centered around zero (output always positive)
- Computationally expensive

$$\sigma(x) = \frac{1}{1+e^{-x}}$$

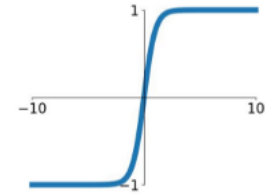


```
# Pytorch implementation  
nn.Sigmoid()
```

Tanh

- No gradient when saturated
- Centered around zero
- Computationally expensive

$\tanh(x)$



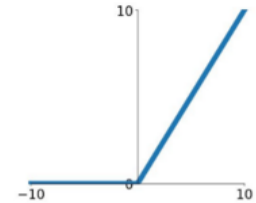
```
# Pytorch implementation  
nn.Tanh()
```

Rectified Linear Unit (ReLU)

- No gradient when saturated (for negative values)
- Not centered around zero (output always positive or zero)
- Computationally & memory cheap
- Back-propagation

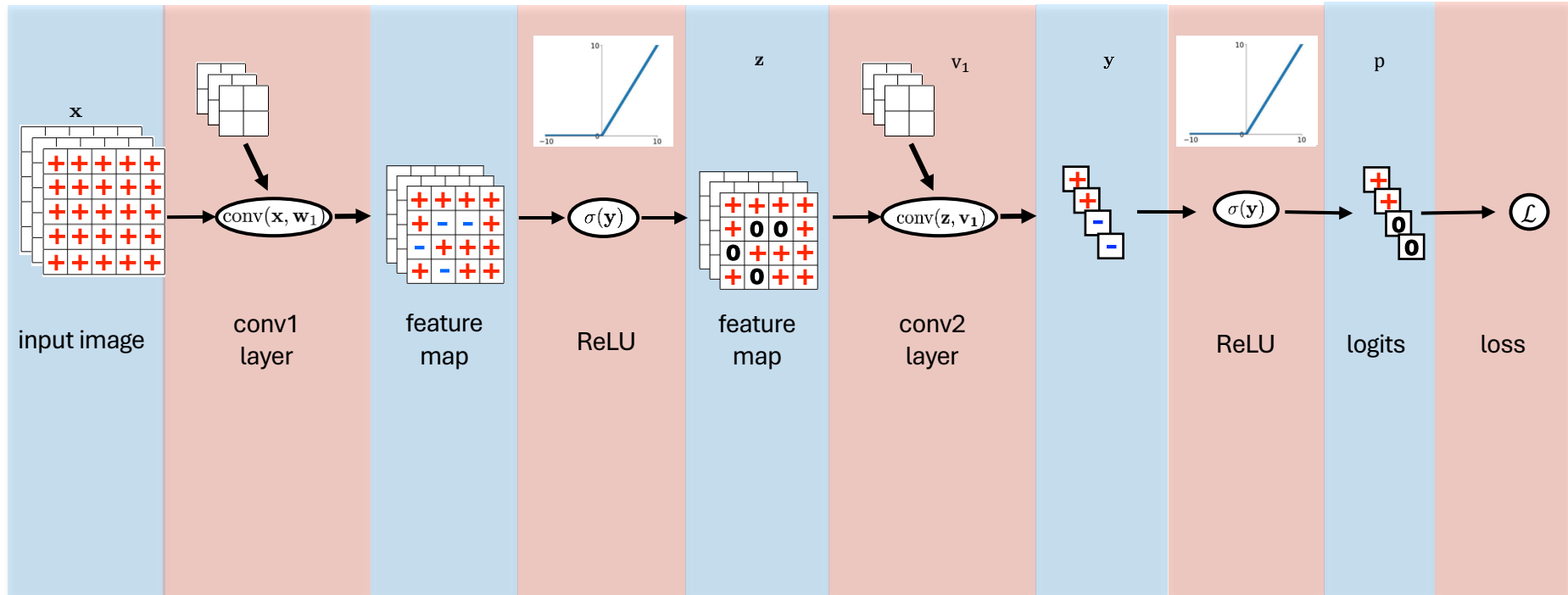
$$\frac{\partial \max(0, x)}{\partial x} = \begin{cases} 0 & x < 0 \\ 1 & \text{otherwise} \end{cases}$$

$$\max(0, x)$$

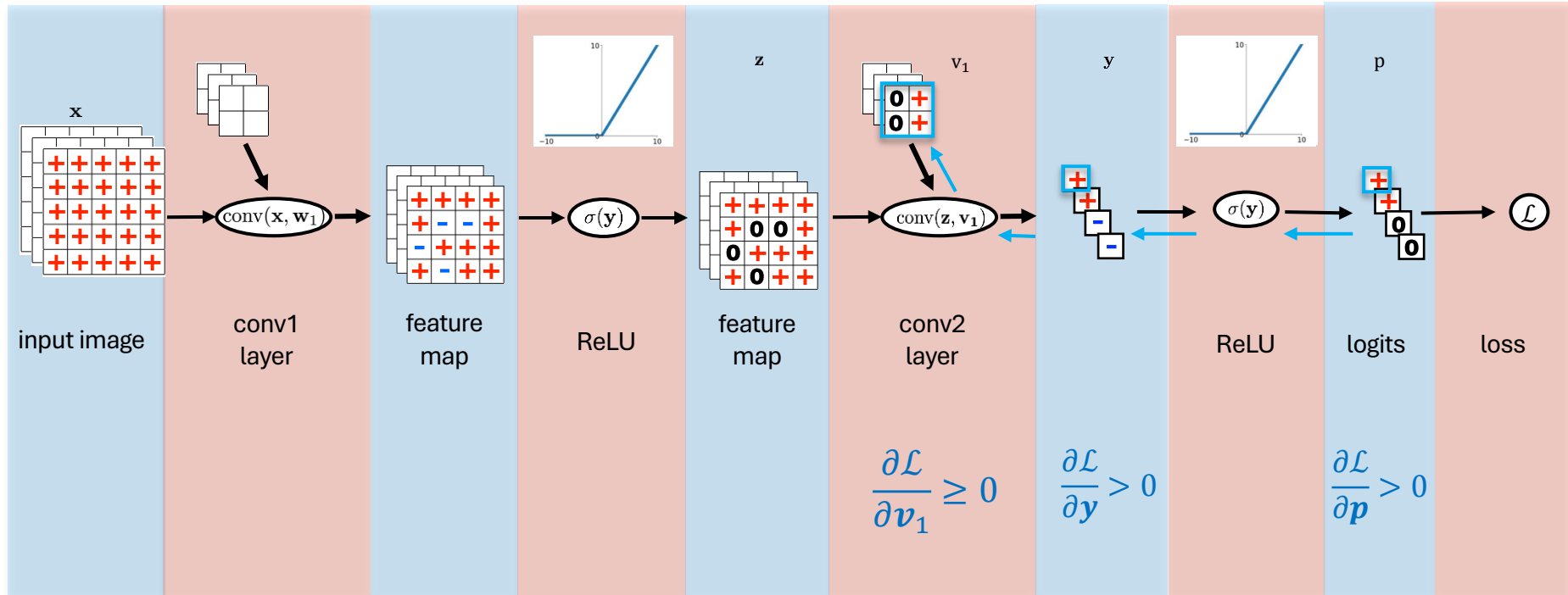


```
# Pytorch implementation  
nn.ReLU(inplace=True)
```

Rectified Linear Unit (ReLU)

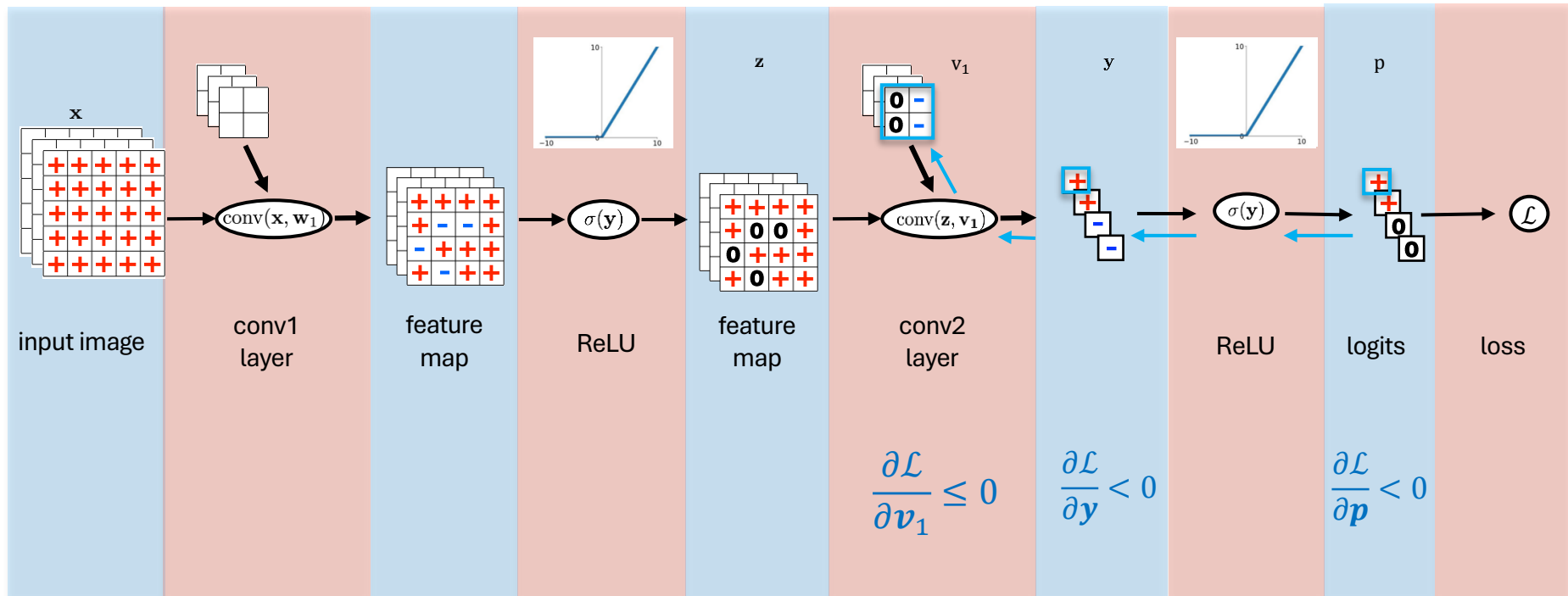


Rectified Linear Unit (ReLU)



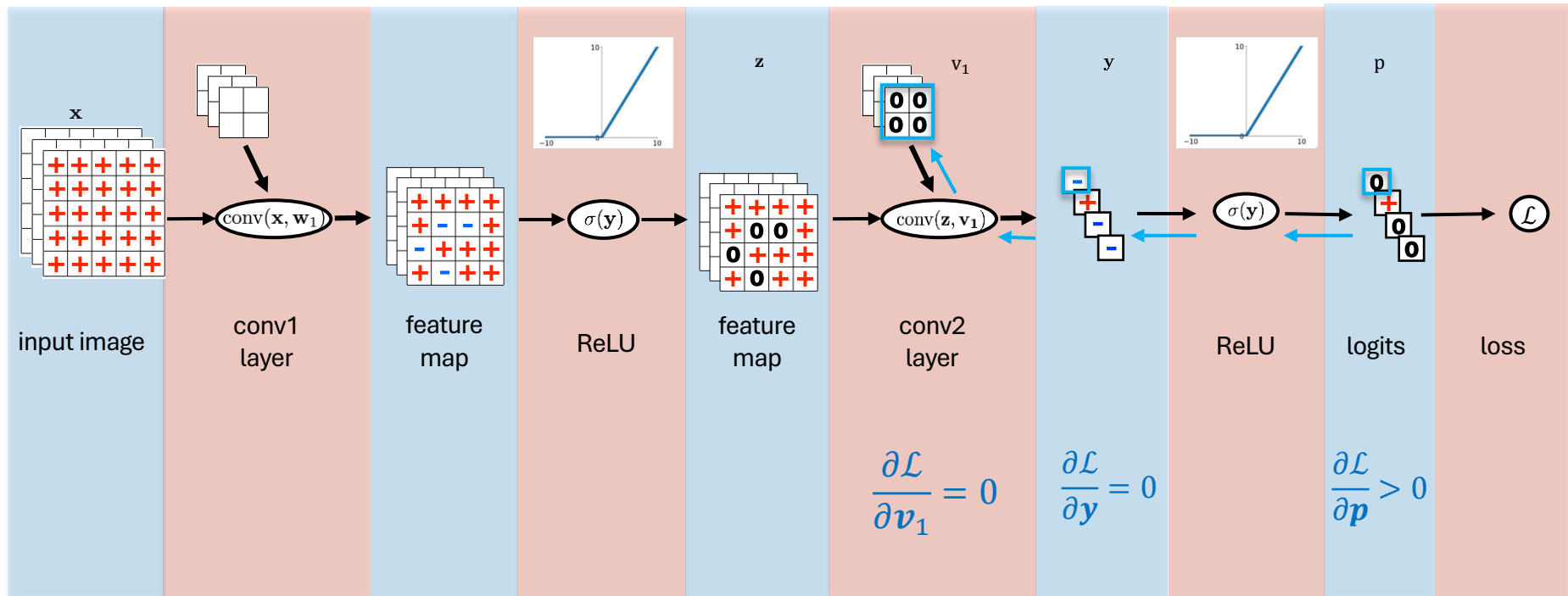
Rectified Linear Unit (ReLU)

- Upstream gradient is positive $\rightarrow$ weights gradient is positive or zero (the given weight remains unchanged)
- Upstream gradient is negative $\rightarrow$ weights gradient is negative or zero



Rectified Linear Unit (ReLU)

- When output of ReLU is zero, gradient propagation stops completely
- “Dead ReLU problem”

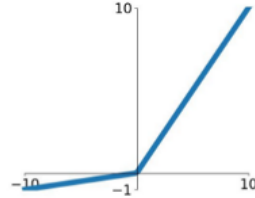


$$\frac{\partial \max(0, x)}{\partial x} = \begin{cases} 0 & x < 0 \\ 1 & \text{otherwise} \end{cases}$$

ReLU alternatives

Leaky ReLU

$\max(0.1x, x)$

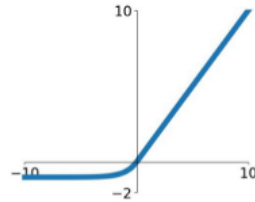


- Does not suffer from dead ReLUs
- Backpropagation:

$$\frac{\partial \max(0.1x, x)}{\partial x} = \begin{cases} 0.1 & x < 0 \\ 1 & \text{otherwise} \end{cases}$$

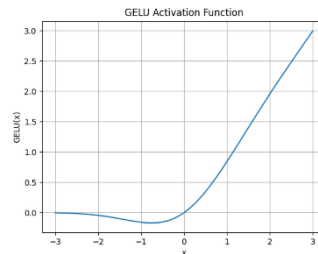
ELU

$$\begin{cases} x & x \geq 0 \\ \alpha(e^x - 1) & x < 0 \end{cases}$$



- Does not suffer from dead ReLUs
- Removes discontinuity around zero

GELU



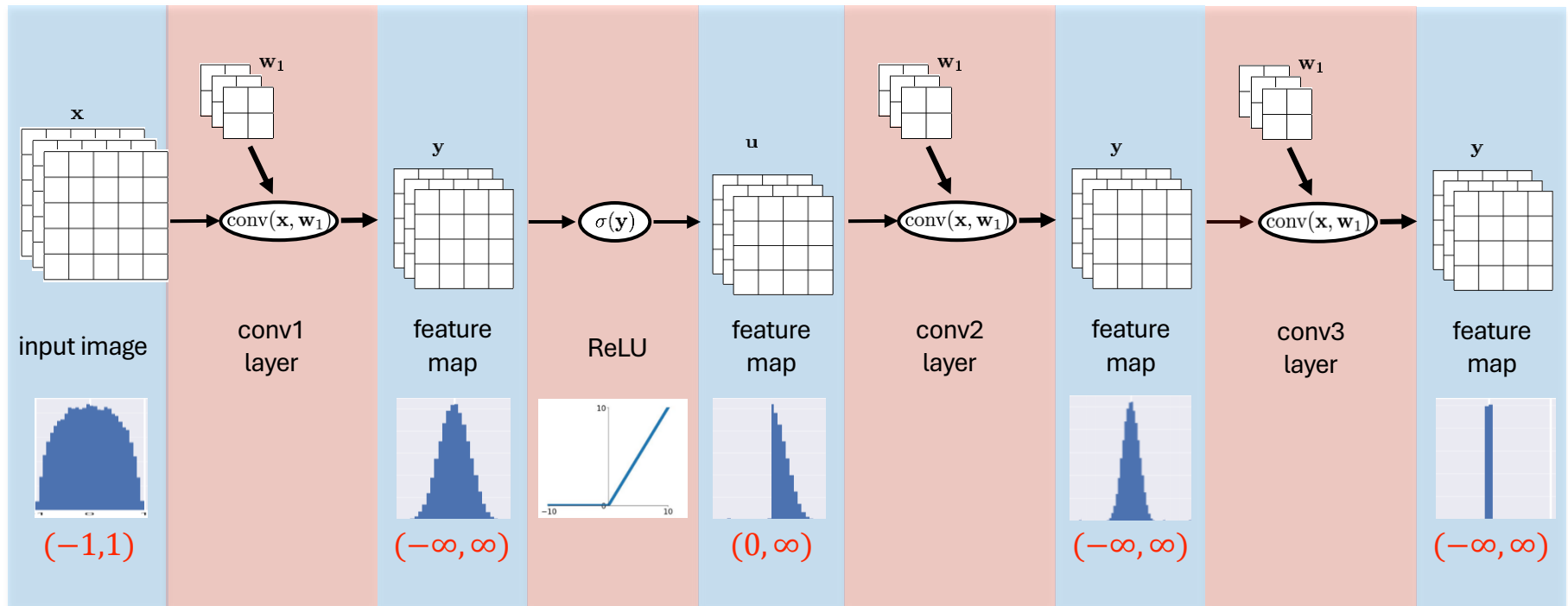
- Does not suffer from dead ReLUs
- Removes discontinuity around zero
- Zero activation for large negative values
- Used in Transformers

Activation Functions

- Many reasonable choices (ReLU, LeakyReLU, ELU, GELU, ...)
- Choosing the best one strongly depends on problem and the data

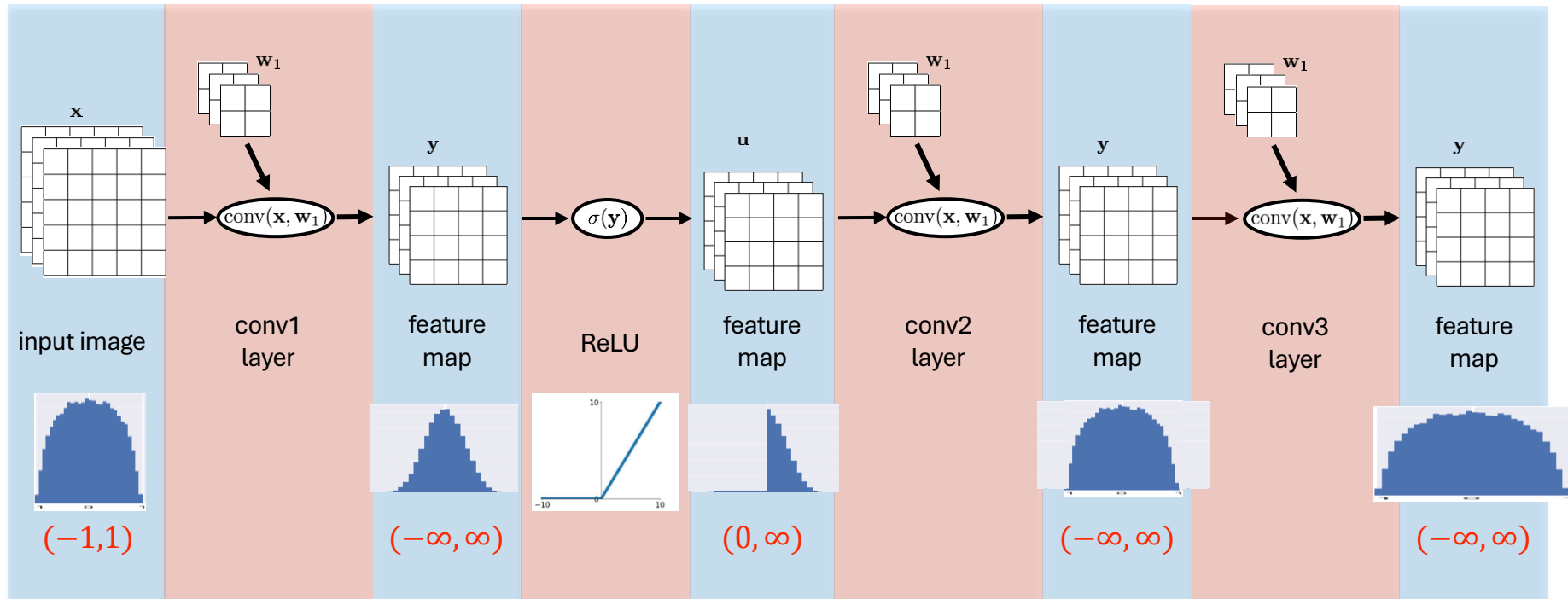
Feature Values

- Small weights cause feature values converge towards zero



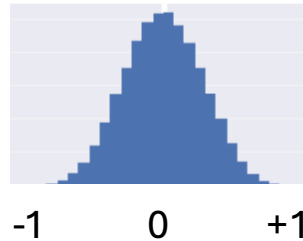
Feature Values

- Small weights cause feature values converge towards zero
- Large weights cause feature values explode towards infinity
- Gradients are also a convolution $\rightarrow$ gradients tend to vanish or explode



Feature Values

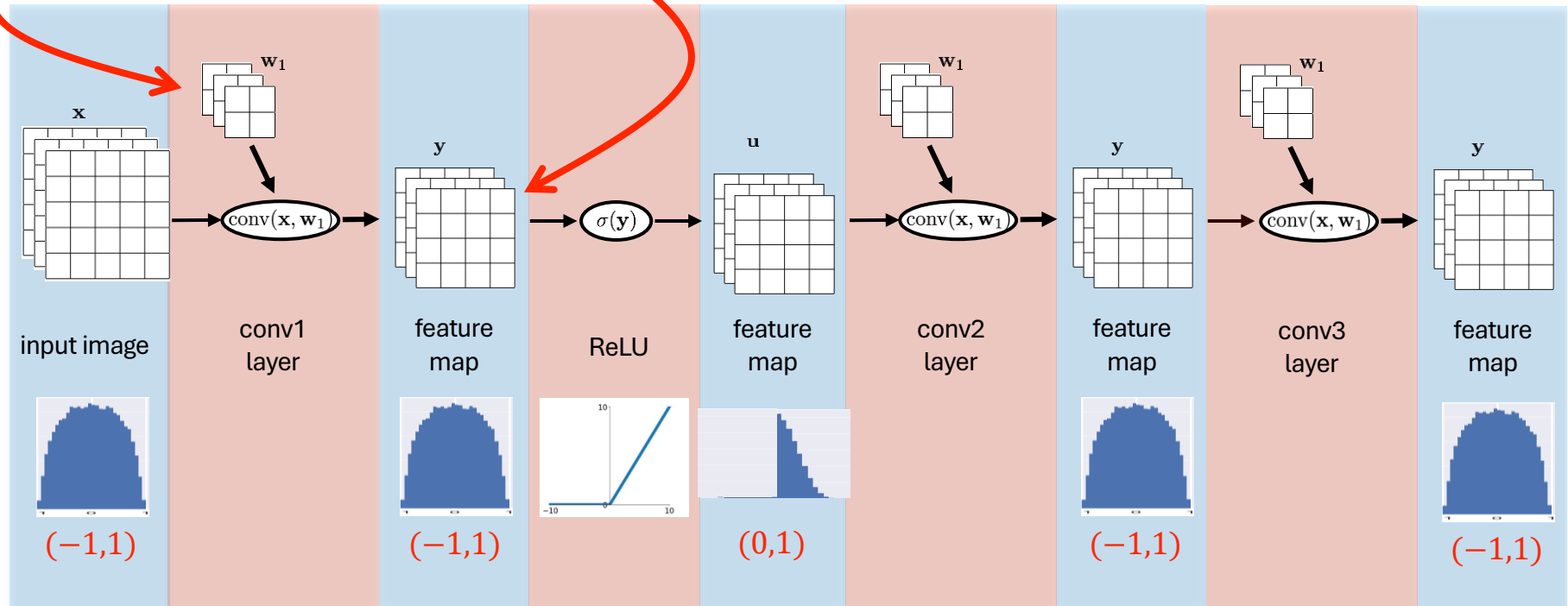
- Values in feature maps can easily explode, diminish or get biased
- Gradients can easily explode or diminish or get biased
- We need to keep reasonable values - what are reasonable values?



$$\mu_x = 0 \quad \text{var}(x) = 1$$

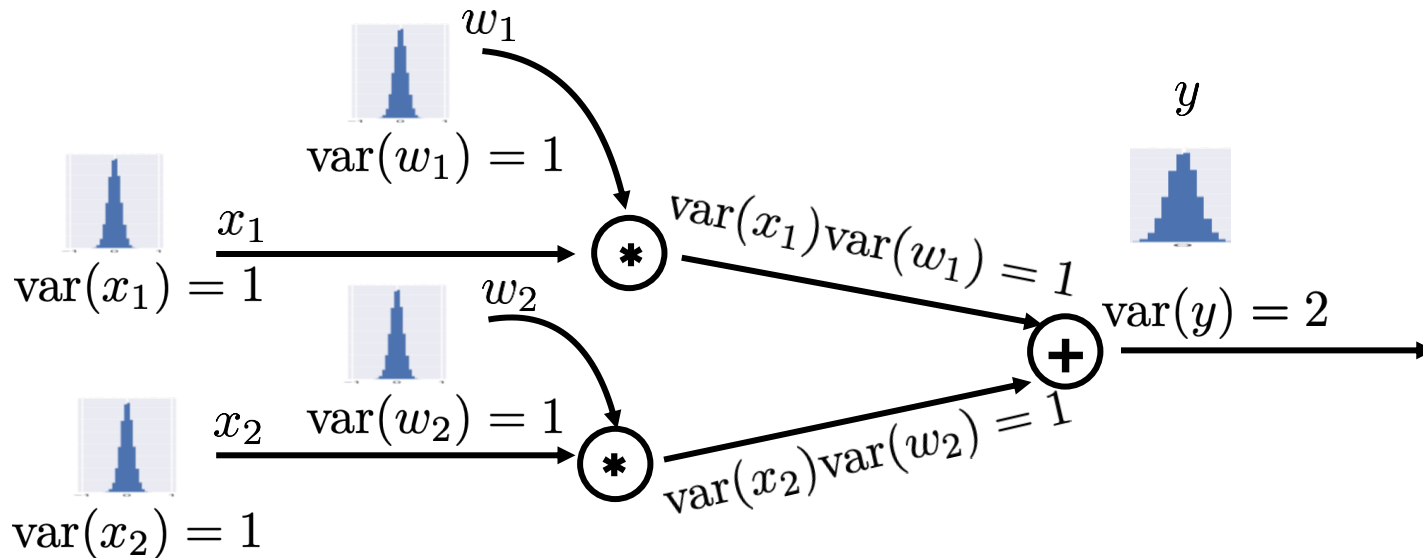
Feature Values

- To keep feature values within “reasonable” range, we need to consider
 1. Weight Initialization
 2. Feature Normalization



Weight Initialization

- How to preserve variance between layers (i.e $\text{var}(y) = \text{var}(x_i)$) ?



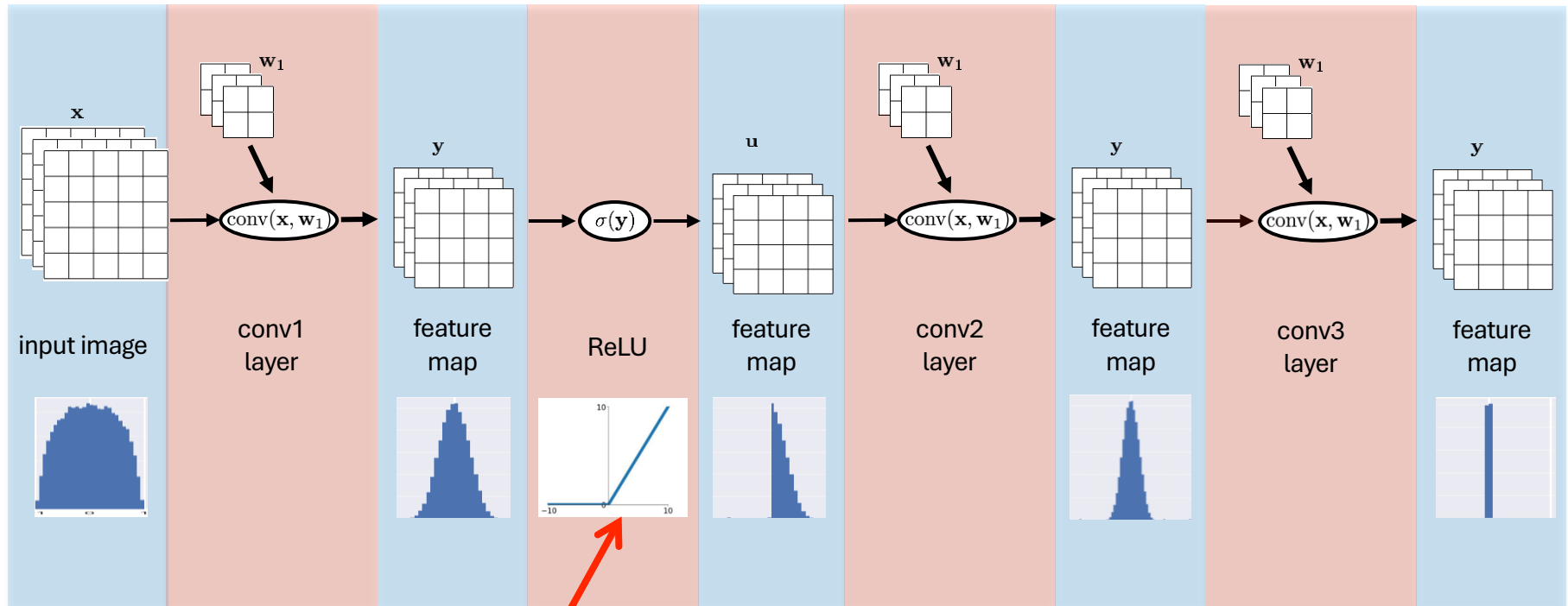
$$\text{var}(x_1 w_1) = (\text{var}(x_1) + \mu_{x_1}^2)(\text{var}(w_1) + \mu_{w_1}^2) - \mu_{x_1}^2 \mu_{w_1}^2 = \text{var}(x_1) \text{var}(w_1) = 1$$

$$\text{var}(y) = \text{var}(x_1 w_1 + x_2 w_2) = \text{var}(x_1 w_1) + \text{var}(x_2 w_2) = 2 \Rightarrow \text{var}(w_i) = \frac{1}{2}$$

$$\text{var}(y) = \text{var}(w_1 x_1 + w_2 x_2 + \dots + w_N x_N) =$$

$$= \sum_{i=1}^N \text{var}(w_i) \text{var}(x_i) \approx N * \text{var}(w_i) \text{var}(x_i) \Rightarrow \text{var}(w_i) = \frac{1}{N}$$

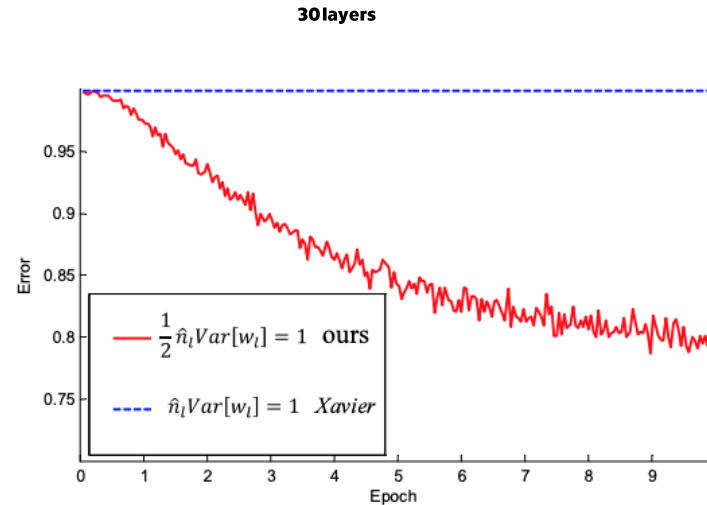
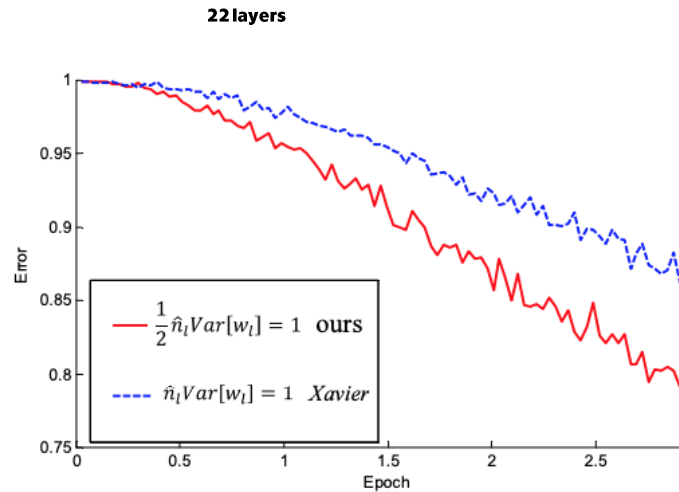
Weight Initialization



What does ReLU do to the variance? It halves it.

Kaiming Initialization

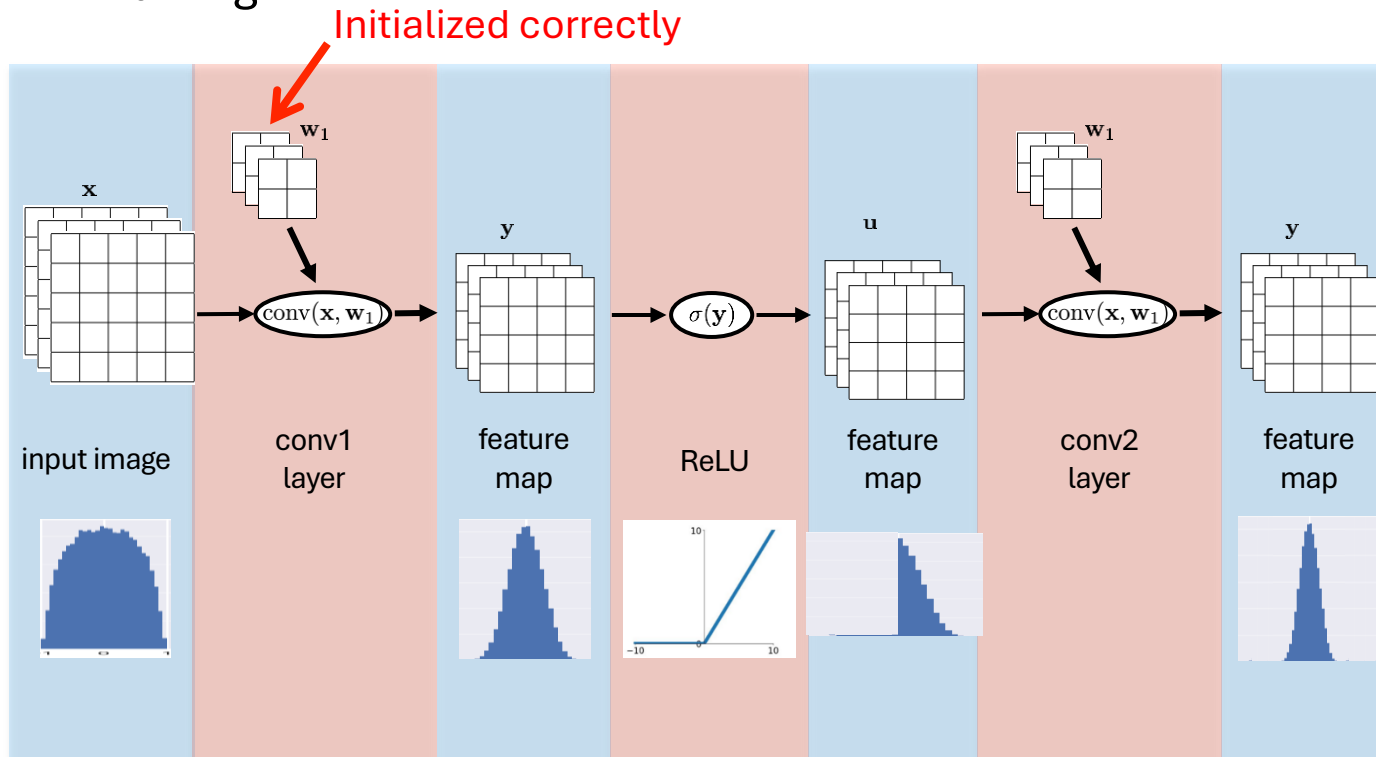
- ReLU reduces variance 2x by itself $\Rightarrow \text{var}(w_i) = \frac{2}{N}$
- Default in PyTorch for Conv layers



He, Kaiming, et al. "Delving deep into rectifiers: Surpassing human-level performance on imagenet classification." ICCV 2015

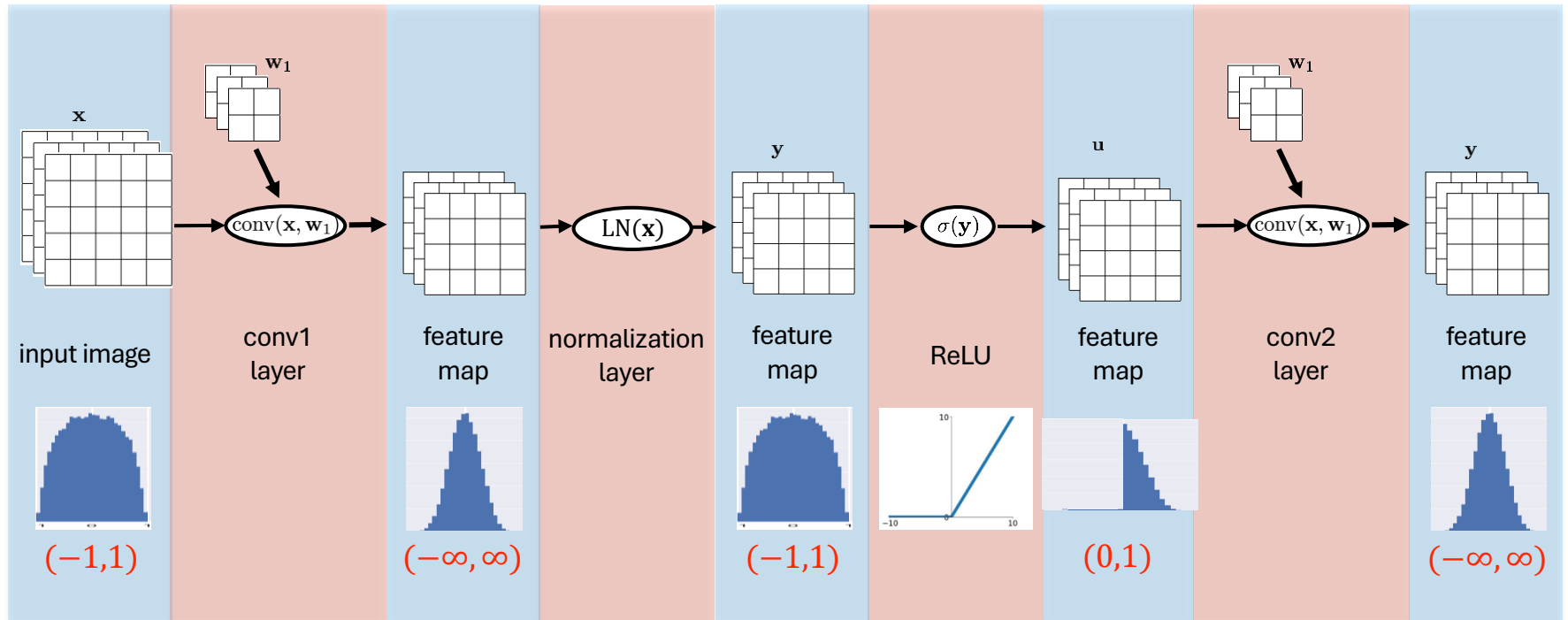
Normalization

- What happens to weights after the first update?
- We need to add feature normalization that keeps everything in range during the training



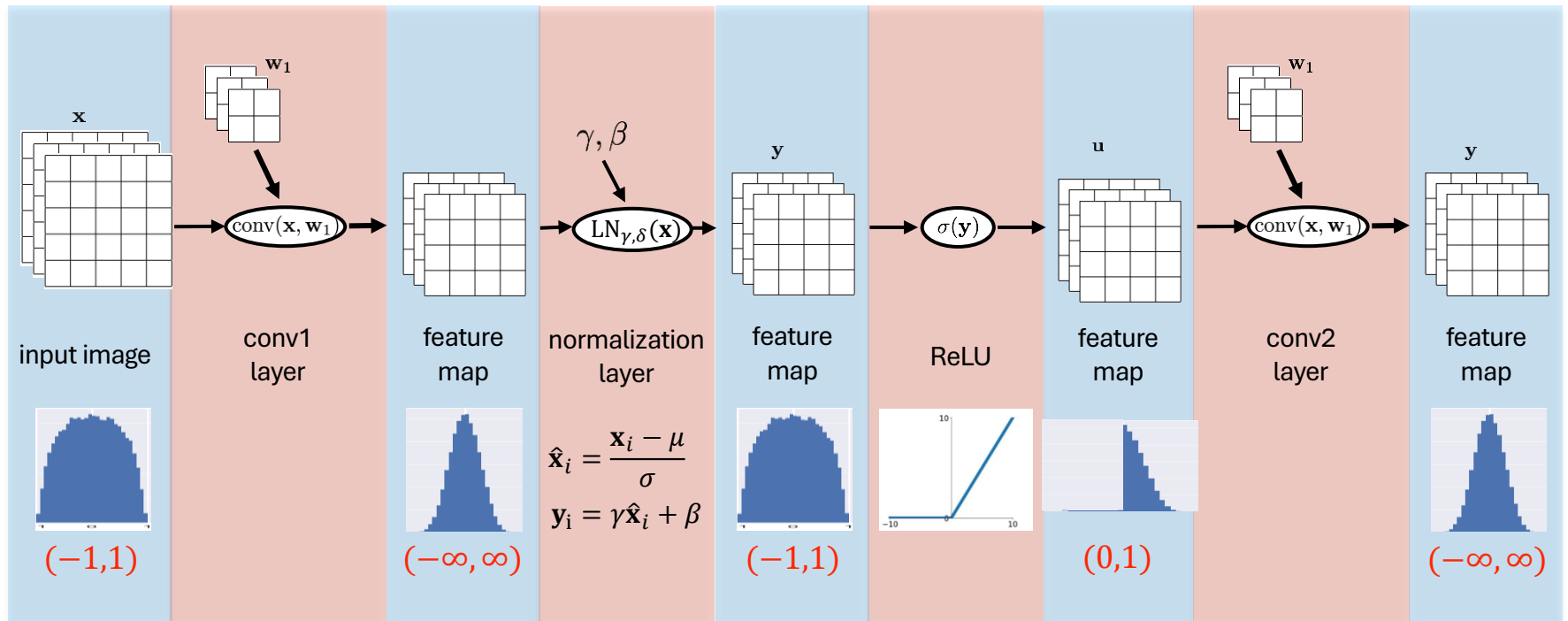
Normalization

- What happens to weights after the first update?
- We need to add feature normalization that keeps everything in range during the training

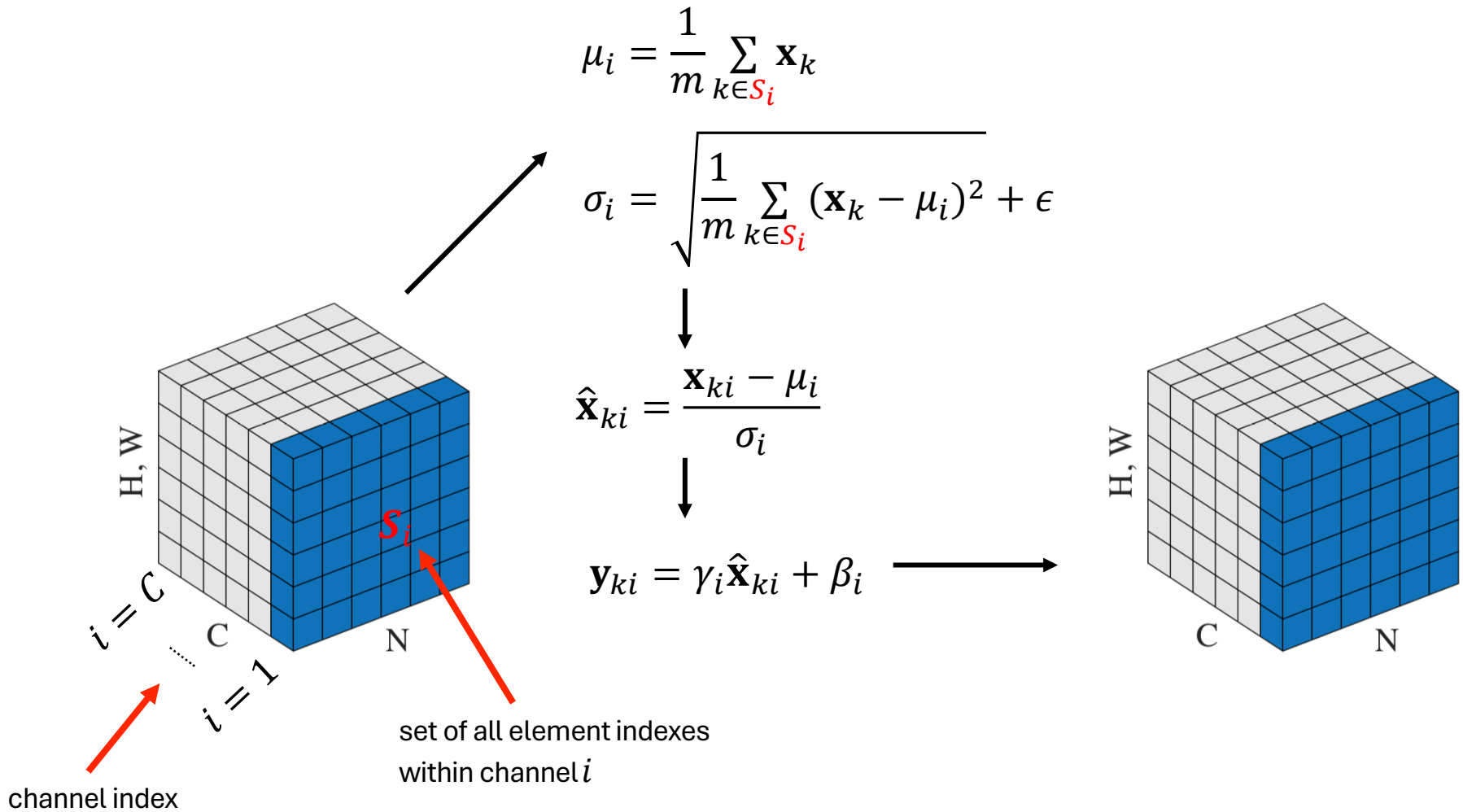


Normalization

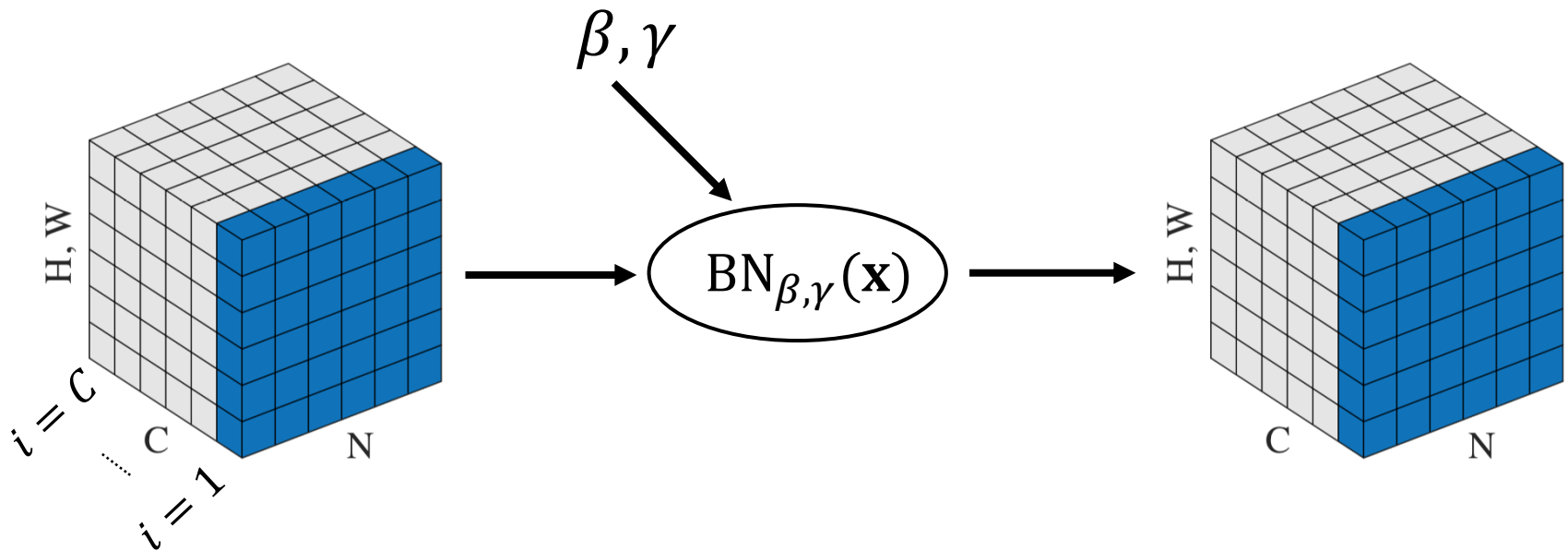
- Subtract mean μ and divide by variance σ to get $\mathcal{N}(0,1)$ distribution
- Learnable affine transformation (γ and β) allows the network to shift the distribution if it's helpful to minimize the loss



Batch Normalization (BatchNorm)



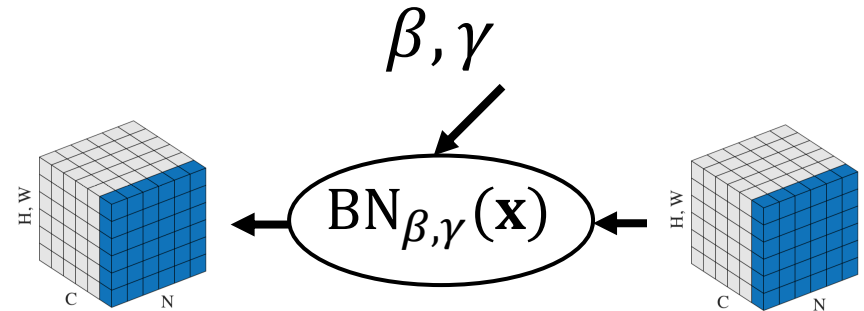
Batch Normalization (BatchNorm)



Batch Normalization (BatchNorm)

```
# Batch x Channels x Width x Height
input = torch.randn(20, 100, 35, 45)

# Batch norm
bn = nn.BatchNorm2d(100)
output = bn(input)
```



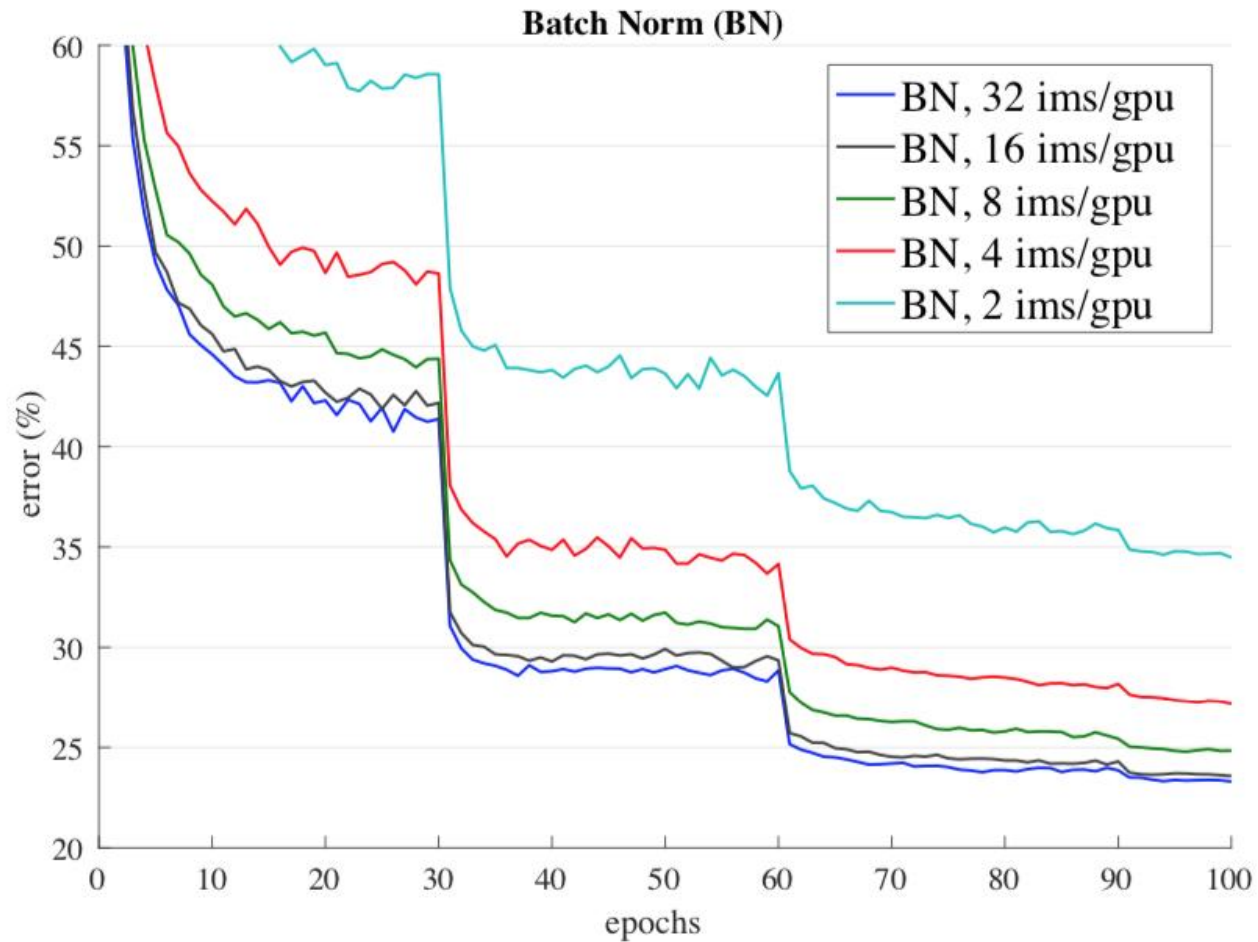
- What is dimensionality of the output?
- The same (20x100x35x45)
- What is dimensionality of mean μ ?
- 100 dimensions
- What is dimensionality of parameter γ ?
- 100 dimensions
- Which μ, σ are used during inference (testing)?
- Estimated in training set, exponential smoothing during inference, ...

Batch Normalization (BatchNorm)

- Why batch normalization helps?
 - It is still unclear, remains an active research topic
- Some claims / hypotheses:
 - Reduces internal covariate shift, i.e. changes in distribution of layer inputs
 - Improves gradient flow and smoother loss
 - Model regularizer: one training example always normalized differently
 - small feature map jittering (dataset augmentation)
 - better generalization

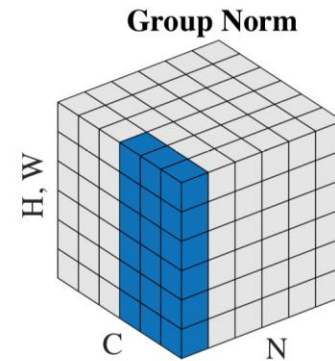
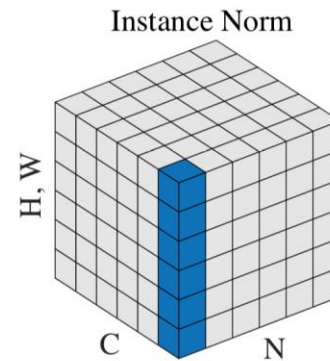
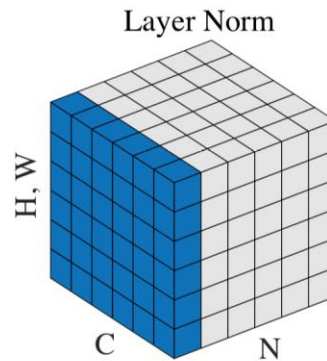
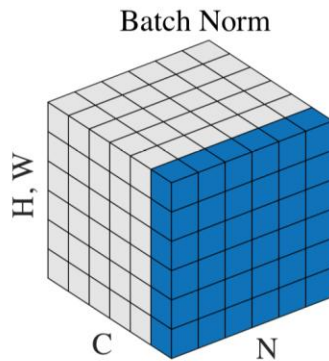
Batch Normalization (BatchNorm)

- **Main drawback:** Does not work well for small batch sizes



Normalization Layers

- Different normalizations are better-suited for different tasks & batch sizes



Best for: Classification

RNN

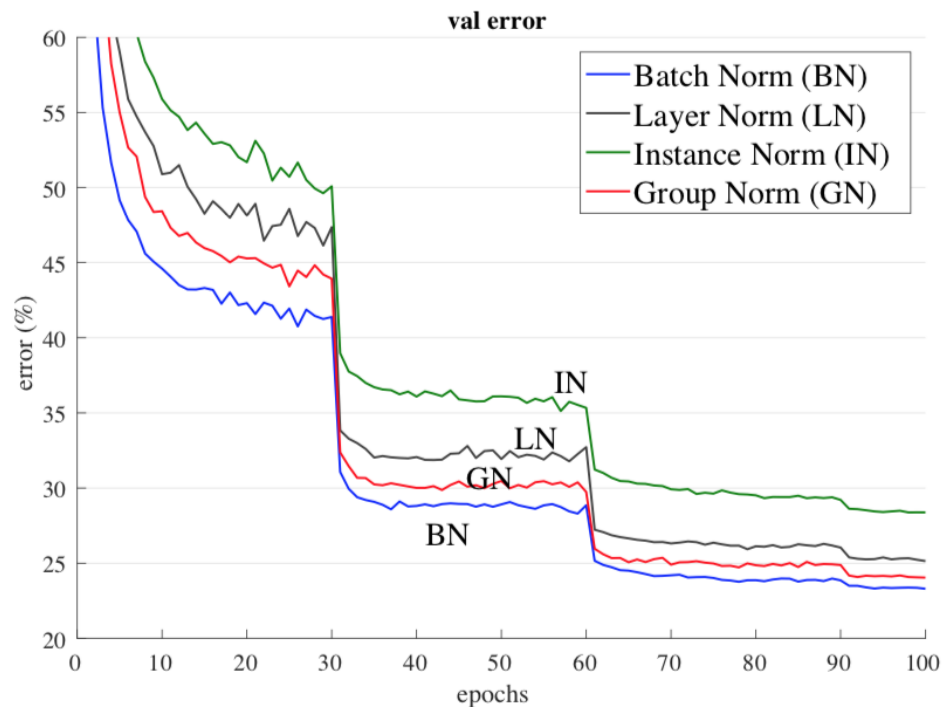
Style transfer

Small batches

Normalization Layers

- Different normalizations are better-suited for different tasks & batch sizes

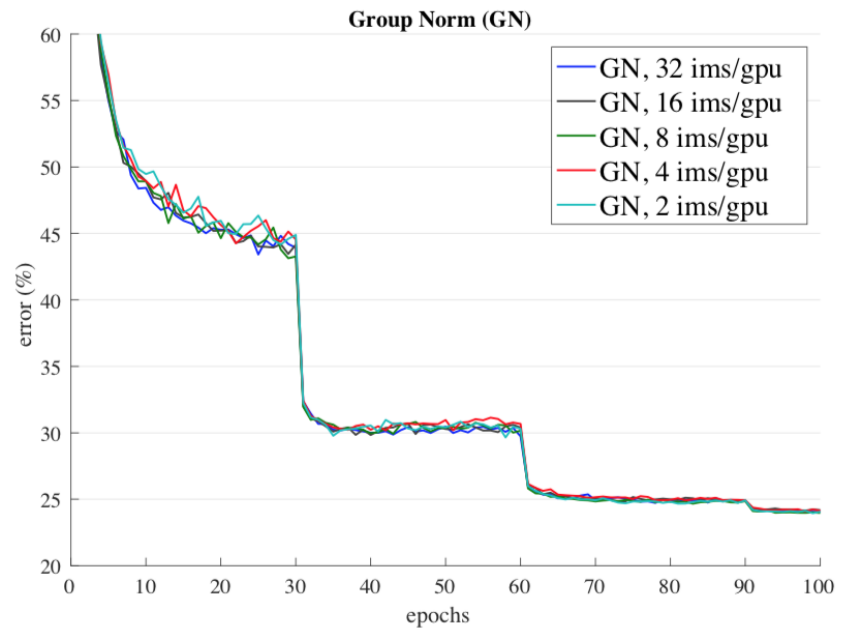
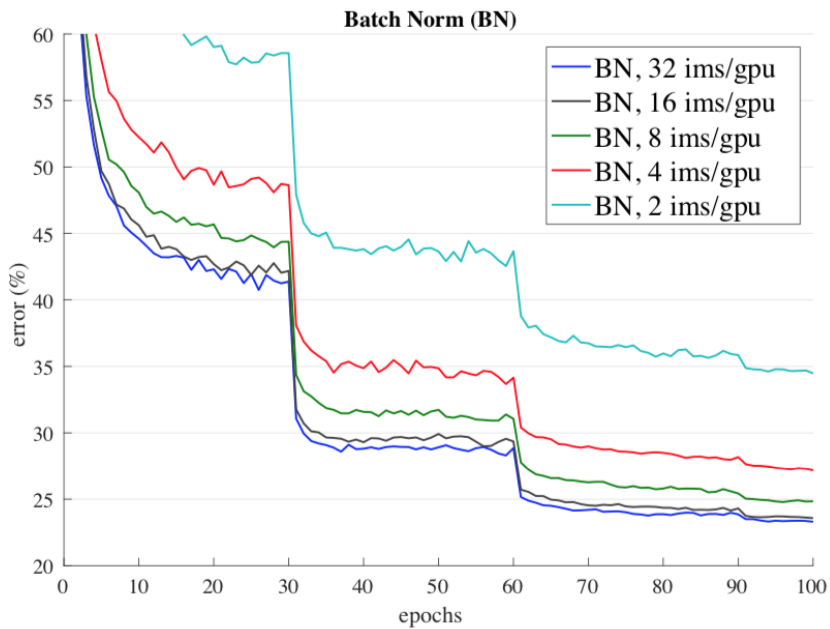
Image classification, batch size = 32



[Wu, He, CVPR, 2018] <https://arxiv.org/pdf/1803.08494.pdf>

Normalization Layers

- Different normalizations are better-suited for different tasks & batch sizes



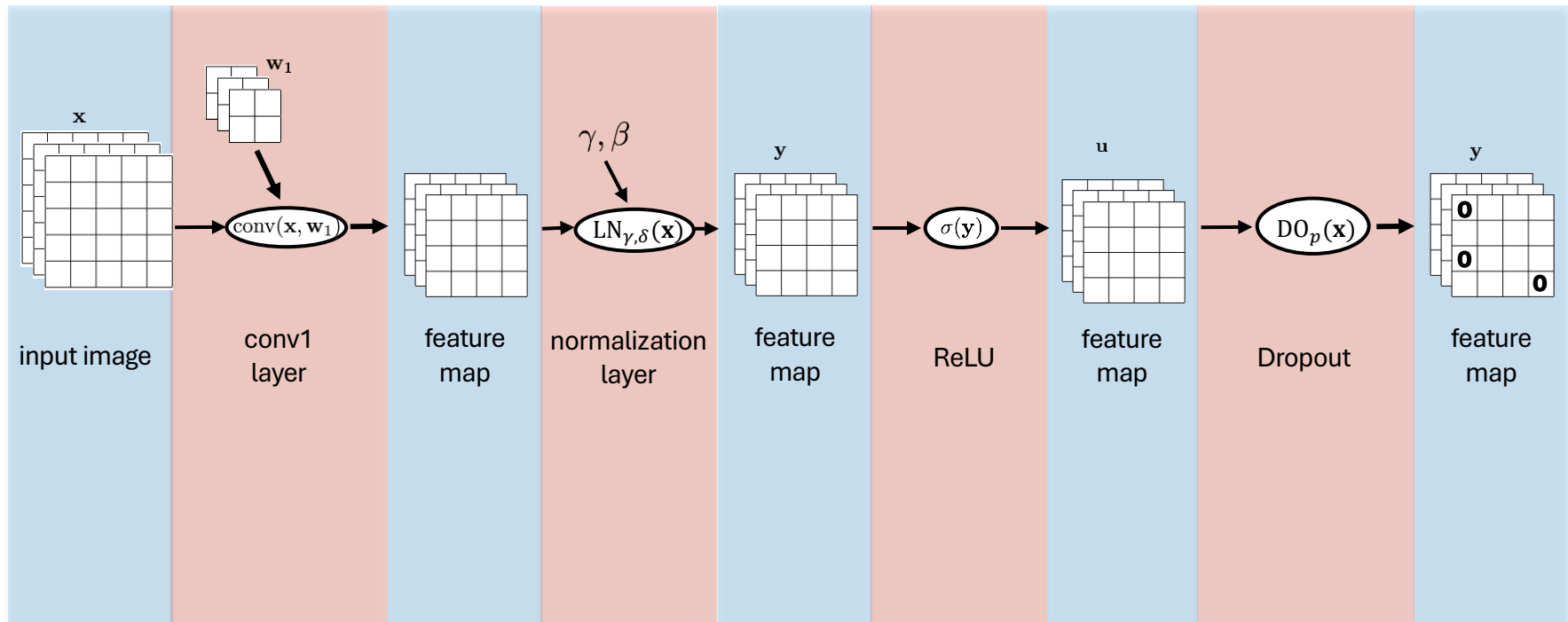
[Wu, He, CVPR, 2018] <https://arxiv.org/pdf/1803.08494.pdf>

Model Regularization

- Prevents overfitting by preferring simple models
- Dropout
- Weight decay

Dropout

- During training, in each iteration randomly set certain fraction of cells to zero
- Reset cells independently, reset random channel (Dropout2d), ...



Dropout

- In Pytorch, is important to switch model between training / eval modes because this controls Dropout layers

```
# Disable dropout -> use all cells
model.eval()

evaluate_model(model)

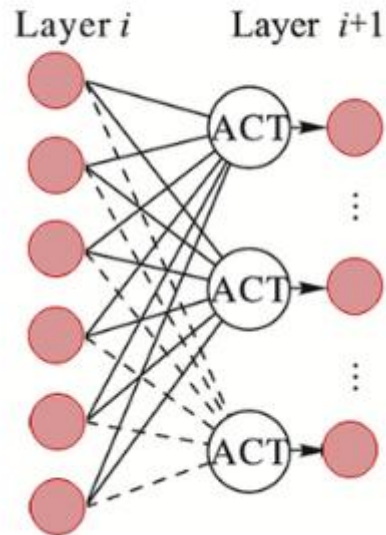
# Re-enable dropout -> randomly zero cells
model.train()

train_one_epoch(model)
```

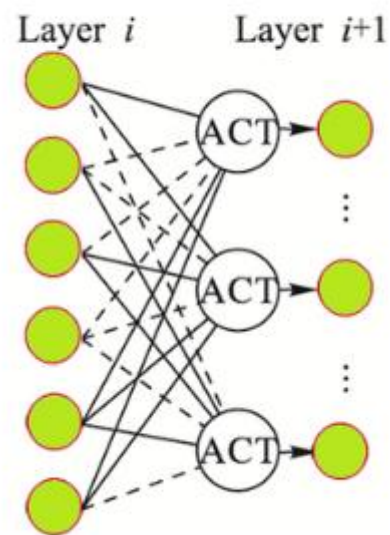
DropConnect

- Randomly drop connections rather than cells

A. Dropout



B. DropConnect



Weight Decay

- We add a L2 norm term to the loss to penalize large weights
- Original loss

$$\mathbf{w}^* = \arg \max_{\mathbf{w}} \left(\prod_i p(y_i | \mathbf{x}_i, \mathbf{w}) \right) = \arg \min_{\mathbf{w}} \mathcal{L}(f(\mathbf{x}, \mathbf{w}), y)$$

- New loss

$$\mathbf{w}^* = \arg \max_{\mathbf{w}} \left(\prod_i p(y_i | \mathbf{x}_i, \mathbf{w}) p(\mathbf{w}) \right) = \arg \min_{\mathbf{w}} \mathcal{L}(f(\mathbf{x}, \mathbf{w}), y) + \lambda \|\mathbf{w}\|^2$$

Weight Decay

- Can be efficiently implemented as part of the update step

$$\mathbf{w}^* = \arg \max_{\mathbf{w}} \left(\prod_i p(y_i | \mathbf{x}_i, \mathbf{w}) p(\mathbf{w}) \right) = \operatorname{argmin}_{\mathbf{w}} \mathcal{L}(f(\mathbf{x}, \mathbf{w}), y) + \lambda \|\mathbf{w}\|^2$$

$$\mathbf{w} := \mathbf{w} - \alpha \left[\frac{\partial \mathcal{L}(f(\mathbf{x}, \mathbf{w}), y)}{\partial \mathbf{w}} + 2\lambda \mathbf{w} \right] = (1 - 2\lambda\alpha) \mathbf{w} - \alpha \frac{\partial \mathcal{L}(f(\mathbf{x}, \mathbf{w}), y)}{\partial \mathbf{w}}$$

```
# Set weight decay as part of the optimizer
optimizer = optim.SGD(model.parameters(), lr=0.01, weight_decay=0.0001)
```

Conclusion

- **Initialization:** set weights to have “normal” values in all layers
- **Normalization:**
 - **BN is reparametrization** of the original NN which has the same expressive power.
 - **BN yields** less sensitivity to vanishing or exploding gradient
 - **BN is model regularizer:** one training example always normalized differently => small feature map jittering (dataset augmentation) => better generalization
 - **BN works well on classification** problems with **larger batches** (>4).
 - Alternatives **LN**, **GN**, **IN** are typically used for **smaller batches, recurrent, generative** and **transformers** networks
- **Regularization:**
 - **Drop-out** reduces overfitting by randomly deactivating neurons/channels
 - **Weight decay** reduces overfitting by reducing reliance on specific weights (kernels)

Competencies gained for the test

- Activation functions (Sigmoid, ReLU, LeakyReLU)
- Weight initialization
- Normalization layers
- Weight decay, dropout