



CS395T: Foundations of Machine Learning for Systems Researchers

Fall 2025

Lecture 2: Gradient Computation in Abstract Neural Networks



Optimization of DNNs: Turing Award citations (2018)

Backpropagation: In a 1986 paper, “Learning Internal Representations by Error Propagation,” co-authored with David Rumelhart and Ronald Williams, Hinton demonstrated that the backpropagation algorithm allowed neural nets to discover their own internal representations of data, making it possible to use neural nets to solve problems that had previously been thought to be beyond their reach. The backpropagation algorithm is standard in most neural networks today.

Improving backpropagation algorithms: LeCun proposed an early version of the backpropagation algorithm (backprop) and gave a clean derivation of it based on variational principles. His work to speed up backpropagation algorithms included describing two simple methods to accelerate learning time.



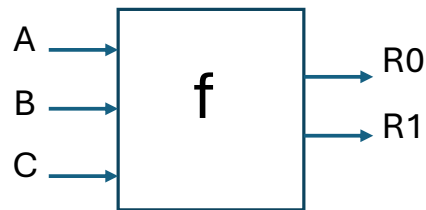
Goal of this lecture

- Standard presentations of back propagation in neural networks
 - Biological metaphors like neurons and synapses come in the way
 - Properties of activation functions are distraction
 - Chain rule of calculus: not obvious how to use it for complicated networks
- Presentation in this lecture
 - Abstraction for neural networks: *parameterized programs*
 - Gradient computation
 - Abstraction (what?): sum over paths
 - Implementation (how?): compositional algorithm on dataflow graph representation
 - Handles complicated neural networks
 - Nonlinear functions like $\sin(x)$, $\cos(x)$, $\sin(\cos(x))$, $e^{\cos(x)}$
 - Arbitrary DAGs including skip connections and weight sharing

Organization

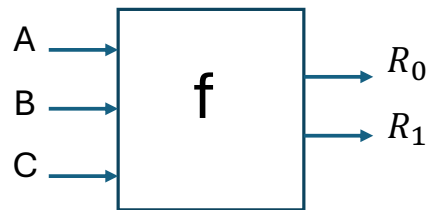
- **Basics:**
 - Partial derivatives, gradients, Jacobians
 - Linear regression:
 - > Parameter optimization: solving linear systems
 - Neural networks: parameter optimization requires solving non-linear systems
 - > Gradient descent
- **Parameterized programs:** abstraction for neural networks
 - Dataflow graph representation
 - Function evaluation: forward propagation
 - Computing gradients: back propagation
 - Parameter optimization using gradient descent

Basic concepts: Partial derivatives, gradients, Jacobians (I)



- Convention:
 - variables are written as capital letters ($A, B, R0$)
 - variable values are written as corresponding small letters ($a, b, r0$)
- Partial derivative of $R0$ wrt A : $\frac{\partial R0}{\partial A}$
 - Function that tells you how much $R0$ changes if A is changed a small amount
 - Value of partial derivative depends on values of A, B, C (consider $f(x) = x^2$)
 - Notation: $\frac{\partial R0}{\partial A}(a, b, c)$ is “value of partial derivative of $R0$ wrt A when input is (a, b, c) ”
- If value of A changes from a to $(a + \Delta a)$, value of $R0$ changes by $\frac{\partial R0}{\partial A}(a, b, c) * \Delta a$

Basic concepts: Partial derivatives, gradients, Jacobians (II)



- **Gradient of R_0 : ∇R_0**

- Column vector containing partial derivatives of R_0

$$\begin{pmatrix} \frac{\partial R_0}{\partial A} \\ \frac{\partial R_0}{\partial B} \\ \frac{\partial R_0}{\partial C} \end{pmatrix}$$

- If input changes from (a, b, c) to $(a + \Delta a, b + \Delta b, c + \Delta c)$, R_0 changes by $(\nabla R_0)^T(a, b, c) * \begin{pmatrix} \Delta a \\ \Delta b \\ \Delta c \end{pmatrix}$

- **Jacobian of f : J_f**

- Matrix whose columns are gradients of outputs

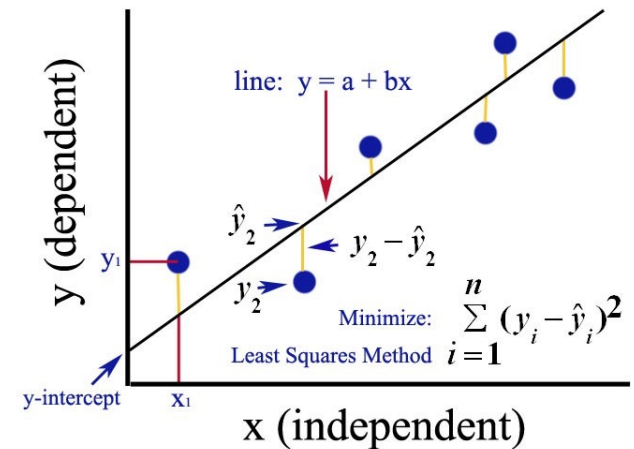
$$J_f = \begin{pmatrix} \frac{\partial R_0}{\partial A} & \frac{\partial R_1}{\partial A} \\ \frac{\partial R_0}{\partial B} & \frac{\partial R_1}{\partial B} \\ \frac{\partial R_0}{\partial C} & \frac{\partial R_1}{\partial C} \end{pmatrix}$$

- If inputs change from (a, b, c) to $(a + \Delta a, b + \Delta b, c + \Delta c)$, outputs change by $J_f^T(a, b, c) * \begin{pmatrix} \Delta a \\ \Delta b \\ \Delta c \end{pmatrix}$

Parameter optimization: linear regression

- Given set of n training data samples $\{(x_i, y_i)\}$, find “best” line $Y = a + b \cdot X$ for given data
 - Model** for training data
 - Model parameters:** a and b
 - Generalization:** model lets you predict y for x not in training sample
- Scoring a proposal (a, b) : **loss function**
 - Compute total square error/residual
 - Detail: we use mean square error (L2 loss)

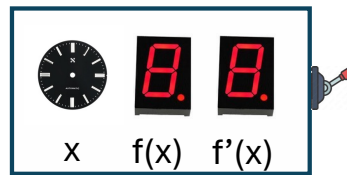
$$Loss(a, b) = \frac{1}{n} \sum_{i=1}^n (y_i - (a + b \cdot x_i))^2$$
- Parameter values that minimize loss
 - Conceptually, a big search problem
 - Better solution: use gradients
 - Set $\nabla Loss(a, b)$ to zero
 - Solve two **linear** equations
- Neural networks are more complicated
 - Parameter optimization requires solving **non-linear** equations
 - Popular method: **gradient descent**



Karl Friedrich Gauss

$$\begin{aligned} n \cdot a + (\sum_i x_i) \cdot b &= \sum_i y_i \\ (\sum_i x_i) \cdot a + (\sum_i x_i^2) \cdot b &= \sum_i x_i y_i \end{aligned}$$

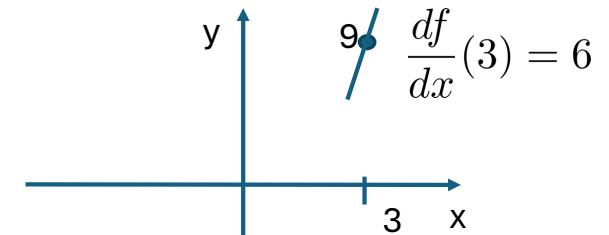
Function minimization using gradient descent (I)



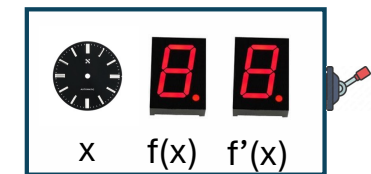
- Problem: given function $f(x)$, find x that minimizes f .
Assumptions:
 - f has minimum and first derivative, but no closed-form expression is available for function or derivative
 - However, we can ask for the value of the function and its derivative at any point
- In general, need iterative search

Function minimization using gradient descent (I)

- Pick some value of x (say 3) and find $f(x)$ (say 9)
 - To see if we have a minimum, find gradient (say 6)
 - Gradient not zero, so not at minimum
 - Good point to try next?



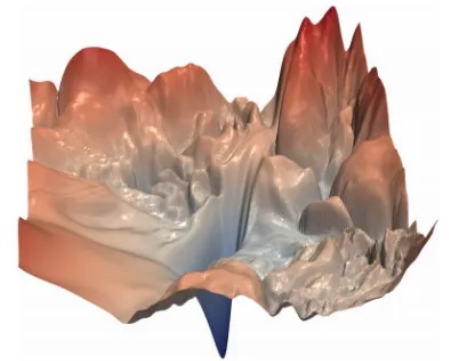
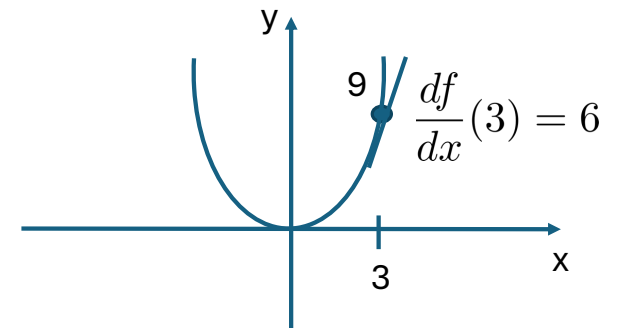
- Gradient gives us two pieces of information
 - Sign:
 - gradient is positive, so increasing x increases function value
 $\Rightarrow$ we should decrease x .
 - Magnitude:
 - if the magnitude of gradient is small, close to the minimum
 $\Rightarrow$ step size should be small
 - if magnitude of gradient is large, may be far from minimum
 $\Rightarrow$ step size should be larger



- This suggests an iterative scheme of the form
$$x_{i+1} = x_i - \frac{df}{dx}(x_i)$$

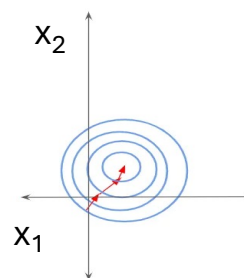
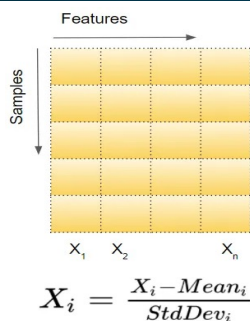
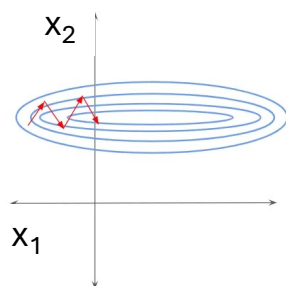
Function minimization using gradient descent (II)

- **Problem:** when magnitude of derivative is large, step size will be big, and we may overshoot minimum
 - Example: function is $y = x^2$ and x_0 is 3.
 - May never converge
- **Solution:** change iterative scheme to $x_{i+1} = x_i - \alpha_i \frac{df}{dx}(x_i)$
 - $0 < \alpha_i < 1$ is a sequence of numbers called *learning rate*
 - Convergence to minimum if sequence satisfies [Robbins-Munro](#) conditions
$$\sum_{i=0}^{\infty} \alpha_n = \infty \quad \sum_{i=0}^{\infty} \alpha_n^2 < \infty$$
 - Safe but slow solution: set α_i to $\frac{1}{i}$
- Iterative scheme for multiple variable function: $\underline{x}_{i+1} = \underline{x}_i - \alpha_i \nabla f(\underline{x}_i)$
- Gradient descent finds local minimum, may not be global minimum



[Hao Li et al.](#)

Variations on theme (I)

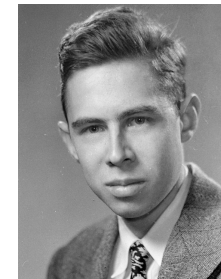
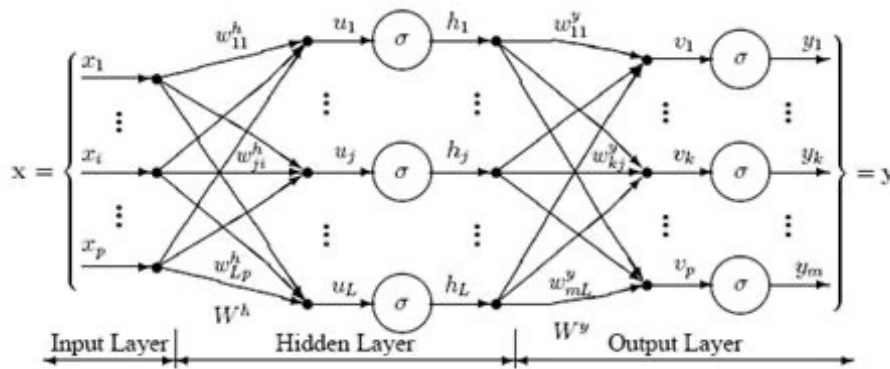


Ketan Doshi

- Rate of convergence
 - How fast is gradient changing?
 - Curvature: second-derivative (Hessian in higher dimensions)
 - Noisy values and gradients: penalty for inaccuracy
- Batch normalization (BN) [Ioffe & Szegedy 2015](#)
 - Reshape optimization landscape to avoid “valleys”
- Other iterative schemes (optimizers): Adagrad, Nesterov momentum, [AdamW](#),
 - <https://www.slideshare.net/SebastianRuder/optimization-for-deep-learning>

Computing Gradients in Programs

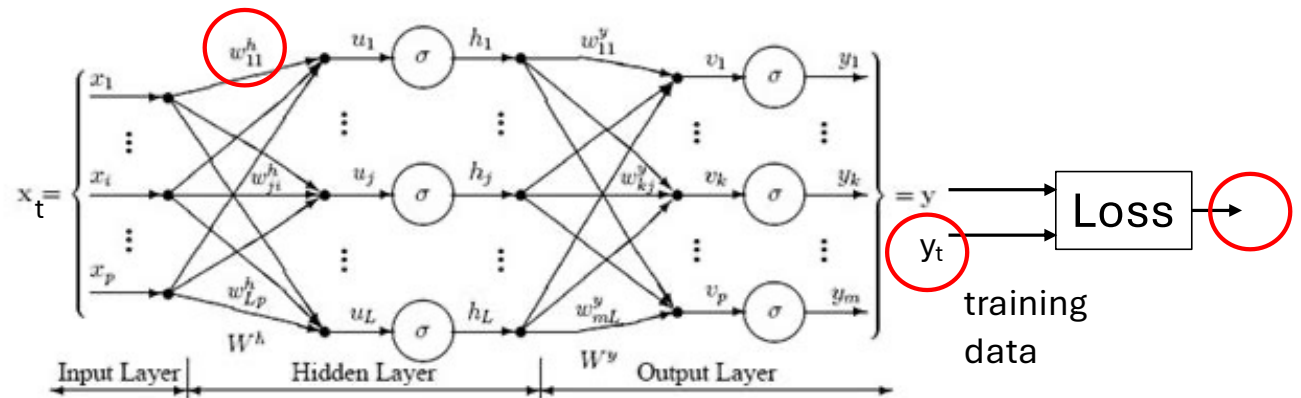
Multilayer Perceptron (MLP) Example



Frank Rosenblatt (Cornell)
Inventor of Perceptron

- **Type:** Inputs: $x_1..x_p$, outputs: $y_1..y_m$ (all $\mathbb{R}$)
- **Scalar view:**
 - Each edge w_{ij} performs a multiplication with a real-valued **parameter**
 - These values are added followed by non-linear operation σ such as tanh, ReLU etc. (known as **activation functions**)
- **Vector view:**
 - Input and output are vectors
 - Each layer performs a dense matrix vector multiplication followed by pointwise (map) non-linear operation σ

Training MLPs



- Optimization problem $f(W; \underline{x}, y)$
 - What values of weights W_{ij} minimize loss for training data?
- Gradient descent
 - Compute derivative of loss w.r.t each weight
 - Use an optimizer from earlier discussion
- Usual presentation of gradient computation: chain rule
- Abstraction of MLP and more complex neural networks: ***parameterized programs***

Parameterized program: running example

- Type of desired function: $\text{real} \times \text{real} \rightarrow \text{real}$
- Training data: set of N 3-tuples $\{(p_i, q_i, t_i)\}$
- Model
 - **Composition** of
 - > **base functions** f_i (may be **nonlinear** such as tanh, sigmoid, sin, cos, ...)
 - > **parameters** W_i : *real*; assume no weight sharing so each weight occurs just once (VGGNet: 138 million parameters)
 - Function written as $R(w; p_i, q_i)$ where w is (w_0, w_1, w_2)
- Parameter optimization
 - Square error for training sample $(p_i, q_i, t_i) = (t_i - R(w; p_i, q_i))^2$
 - Goal: choose (w_0, w_1, w_2) to minimize mean square error (L2 loss)
$$\text{Loss}(w_0, w_1, w_2) = \frac{1}{N} \sum_{i=1}^N (t_i - R(w; p_i, q_i))^2$$

```
Function R(P,Q) {  
    A = W0*P  
    B = f0(A,Q)  
    C = W1*B  
    D = W2*B  
    E = f1(C)  
    F = f2(D)  
    R = f3(E,F)  
    return R  
}
```

Parameter optimization

- Find derivatives of Loss wrt W_0, W_1, W_2 :

$$\text{Loss}(w_0, w_1, w_2) = \frac{1}{N} \sum_{i=1}^N (t_i - R(w; p_i, q_i))^2$$

$$\frac{\partial \text{Loss}}{\partial W_0}(w_0, w_1, w_2) = -\frac{2}{N} \sum_{i=1}^N (t_i - R(w; p_i, q_i)) \frac{\partial R}{\partial W_0}(w; p_i, q_i)$$

$$\frac{\partial \text{Loss}}{\partial W_1}(w_0, w_1, w_2) = -\frac{2}{N} \sum_{i=1}^N (t_i - R(w; p_i, q_i)) \frac{\partial R}{\partial W_1}(w; p_i, q_i)$$

$$\frac{\partial \text{Loss}}{\partial W_2}(w_0, w_1, w_2) = -\frac{2}{N} \sum_{i=1}^N (t_i - R(w; p_i, q_i)) \frac{\partial R}{\partial W_2}(w; p_i, q_i)$$

$$\nabla_W R(w; p_i, q_i)$$

Derivatives are complex, non-linear functions
so use gradient-descent for parameter optimization

Function $R(P, Q)$ {

$A = W_0 * P$

$B = f_0(A, Q)$

$C = W_1 * B$

$D = W_2 * B$

$E = f_1(C)$

$F = f_2(D)$

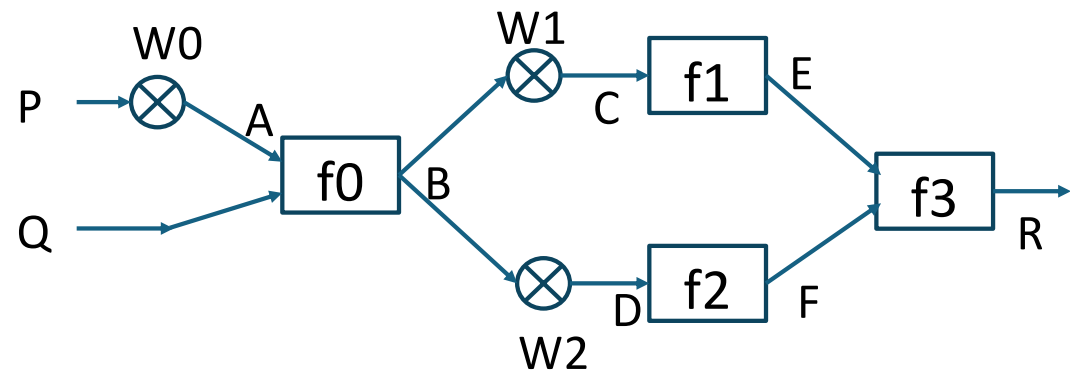
$R = f_3(E, F)$

return R }

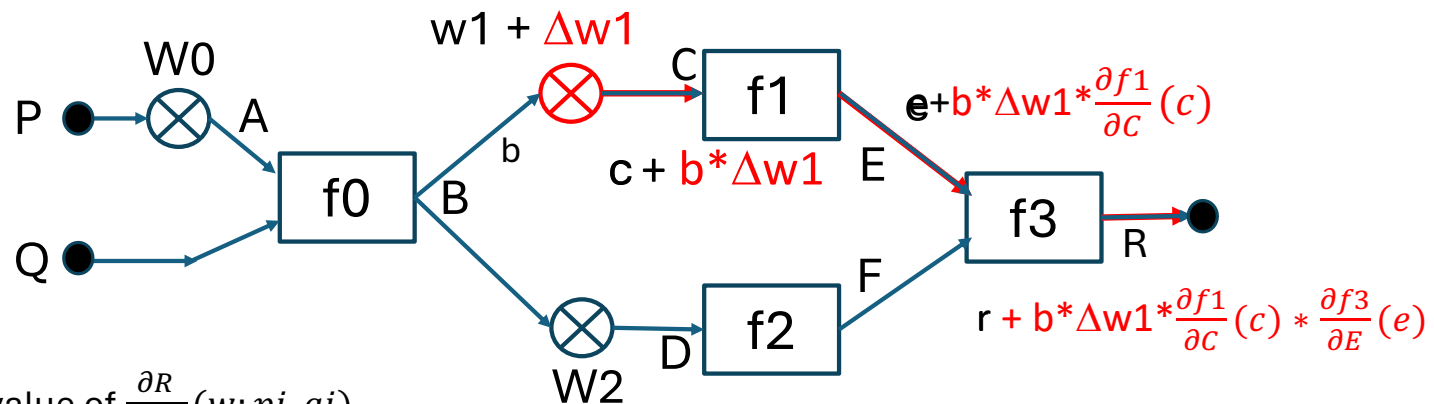
Parameterized program as flow graph

- Useful to represent multiplication by weights differently from functions f_i
 - Functions f_i are fixed but weights change during training
- Execution (**forward propagation**)
 - Sequential: execute nodes in any topological order
 - Parallel dataflow: node executes when inputs are available
- All values at intermediate points ($A, B, C, \dots$) are stored
 - Needed for gradient computations

```
Function R(P,Q) {  
  A =  $W_0 * P$   
  B =  $f_0(A, Q)$   
  C =  $W_1 * B$   
  D =  $W_2 * B$   
  E =  $f_1(C)$   
  F =  $f_2(D)$   
  R =  $f_3(E, F)$   
  return R  
}
```

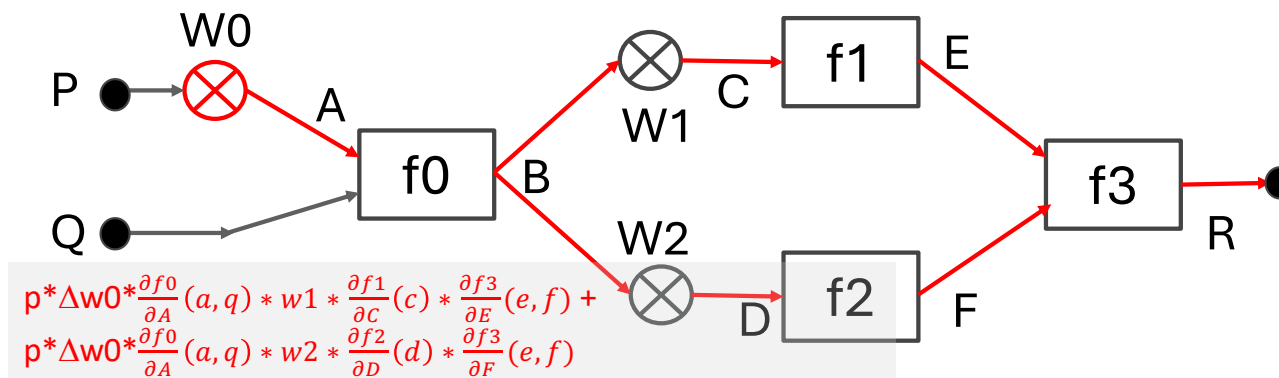


Value of $\frac{\partial R}{\partial W_1}(w; p_i, q_i)$



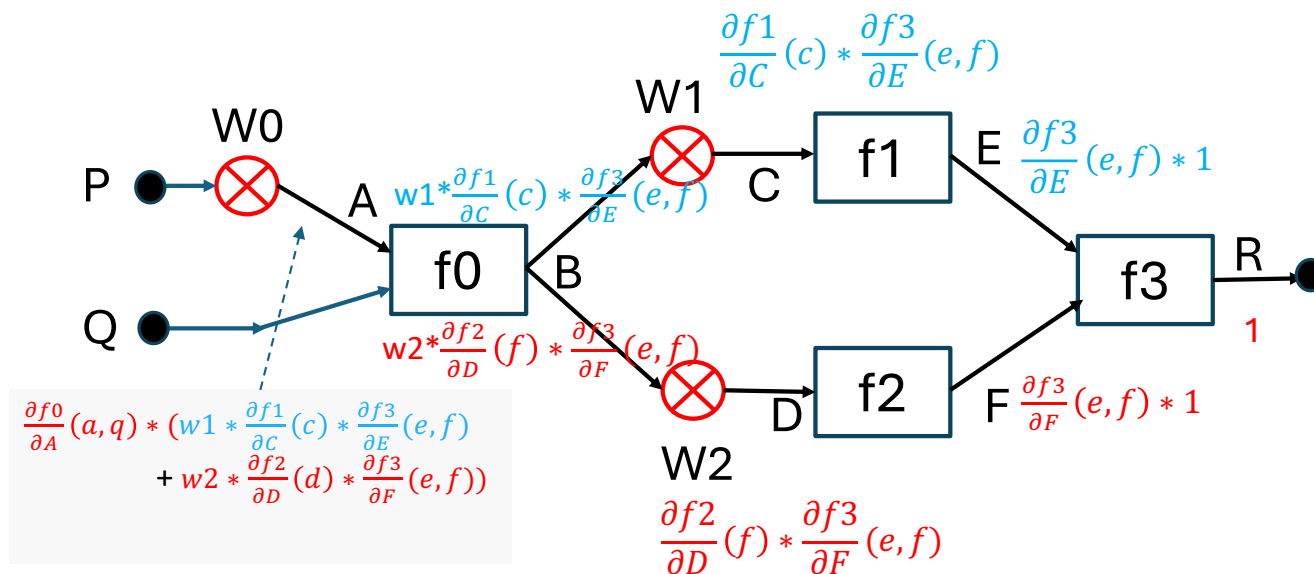
- To find value of $\frac{\partial R}{\partial W_1}(w; p_i, q_i)$
 - Multiply values of partial derivatives of all vertices on path from W_1 to R
 - Multiply result by value of input to W_1 (i.e., b)
 - Result: $b * \frac{\partial f_1}{\partial c}(c) * \frac{\partial f_3}{\partial e}(e)$
 - Compute the product either forwards or backwards along path
- In general, given path $h : X \xrightarrow{*} Y$
 - $\pi(h)$ = product of derivative values of nodes on h excluding X and Y
 - $\pi(h) = 1$ for empty path or if there are no intermediate nodes
- $\frac{\partial R}{\partial W_1}(w; p_i, q_i) = b * \pi(W_1 \xrightarrow{*} R)$

Value of $\frac{\partial R}{\partial w_1}(w; p_i, q_i)$



- In general, there is a DAG from vertex representing weight to the output
 - Example: weight w_0
- Value of partial derivative: **sum over paths from vertex to output**
 - Enumerate all paths from weight to output and add up the contributions of all paths
 - Notation: $H(n)$: set of paths from node n to output
 - Intuition: derivatives make this a linear problem, so *superposition of paths* works
- Problems:
 - Treats DAG like tree so could do exponential computation in size of DAG. More efficient solution?
 - What order should we compute $\frac{\partial R}{\partial(w_0)}(w; p_i, q_i)$, $\frac{\partial R}{\partial(w_1)}(w; p_i, q_i)$ and $\frac{\partial R}{\partial(w_2)}(w; p_i, q_i)$?

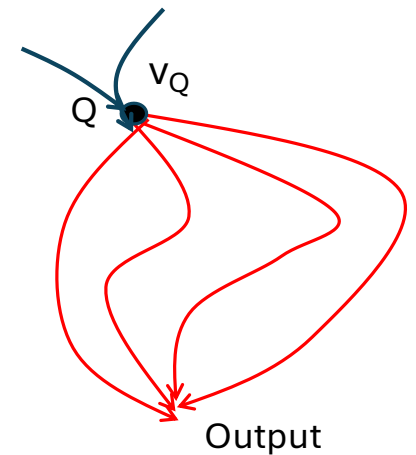
Efficient computation of all derivatives



- Compute derivative of output wrt every variable/edge ($\frac{\partial R}{\partial A}, \frac{\partial R}{\partial C}, \frac{\partial R}{\partial F} \dots$)
 - Real number on each edge
 - Derivatives of output wrt weights can be computed from this
- Traverse DAG in reverse topological order of variables for computation
 - $\frac{\partial R}{\partial R} = 1$
 - **Transfer functions** to propagate derivative from function output to its inputs

Summary of derivatives computation

- At each point Q , compute $v_Q: \mathbb{R}$ where $v_Q = \frac{\partial \text{Output}}{\partial Q}(W; p_i, q_i) = \sum_{h \in H(Q)} \pi(h)$
- Small tweak to handle weight-sharing
- Called back-propagation in ML literature
- Vanishing and exploding gradients
 - Multiplying sequence of small or large numbers



$v_{in} = 1$	$v_{in} = w_i * v_{out} \text{ and } \frac{\partial \text{Output}}{\partial w_i} = v_{out} * x$
$v_{in} = (v1_{out} + v2_{out})$	$v1_{in} = \frac{\partial f}{\partial x}(x, y) * v_{out}$

Back to running example

- Optimization problem: choose W_i values to minimize loss

$$\text{Loss}(w_0, w_1, w_2) = \frac{1}{N} \sum_{i=1}^N (t_i - R(w; p_i, q_i))^2$$

$$\frac{\partial \text{Loss}}{\partial W_0}(w_0, w_1, w_2) = -\frac{2}{N} \sum_{i=1}^N (t_i - R(w; p_i, q_i)) \frac{\partial R}{\partial W_0}(w; p_i, q_i)$$

```

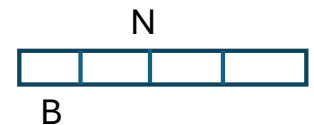
Initialize weights to random values
for #epochs do {
  GradientVector = 0
  for each training sample (pi,qi,ti) do {
    perform forward propagation and compute
    perform backpropagation and compute weight derivatives
    update GradientVector with products}
  scale GradientVector by -2/N
  use GradientVector to update weights using gradient descent step
}
    
```

Function $R(P,Q)$ {
 $A = W0 * P$
 $B = f0(A,Q)$
 $C = W1 * B$
 $D = W2 * B$
 $E = f1(C)$
 $F = f2(D)$
 $R = f3(E,F)$
 return R }

Variations on theme (II)

– How many training data samples should we use before updating weights?

- (Batch) Gradient-descent: use entire training data set to compute gradient.
 - Disadvantage: learn only after all training data has been processed.
- (Mini-)Batching: divide training data into subsets of size B , update weights after each subset is processed and use as initial weights for next subset.
 - Disadvantage: choosing B ? 50-256 is common.
 - Useful to randomize training data set
- Stochastic gradient-descent (SGD): $B = 1$. Can be noisier than previous approaches.



– Initialization: we assumed random initialization

- Can also exploit prior (domain knowledge)
- He initialization
- Xavier (Glorot) initialization
- Good discussion: <https://arxiv.org/abs/1704.08863>

Variations on theme (III)

Other loss functions

- Regularization: add a term that penalizes weights that are not “desirable”

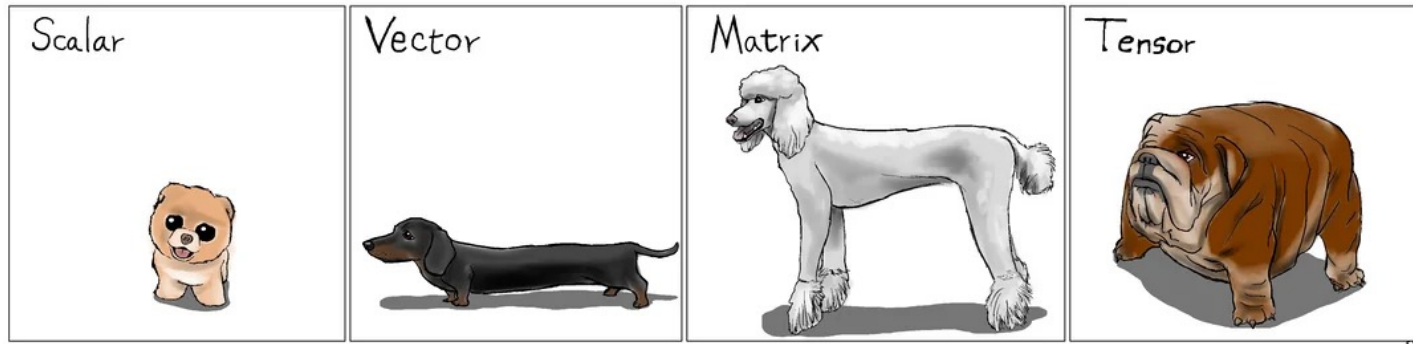
$$\text{Example: } \text{Loss}(w_0, w_1, w_2) = \frac{1}{N} \sum_{i=1}^N (t_i - R(W; p_i, q_i))^2 + \lambda \|W\|^2$$

Called *ridge regularization*: penalizes large weight values

Many other forms of regularization

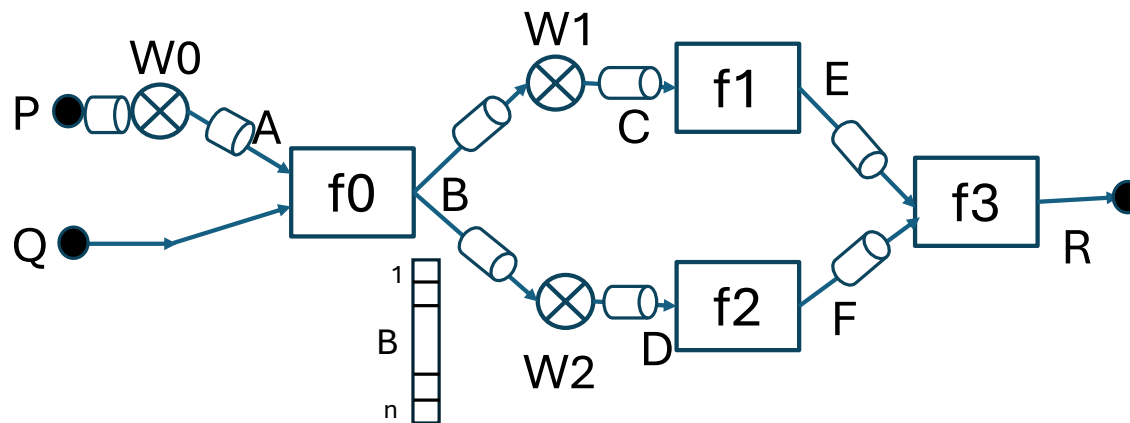
- Loss for classification problems
 - KL-divergence, cross-entropy loss (see later)

Generalization to vectors, matrices, tensors



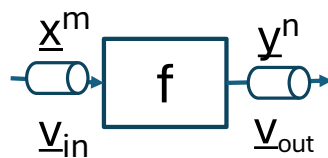
How to understand data with dogs
Karl Stratos ([Reddit post](#))

Generalization to vectors (matrices, tensors are similar)



- Programs
 - Inputs can be scalars or **vectors**, but **output is still a scalar** (loss)
 - Weight matrices (need not be square)
 - Function type: vector $\rightarrow$ vector
- Compute gradients instead of derivatives
 - Variable is vector of size $n \rightarrow$ gradient of output wrt this vector is also vector of size n
 - Intuition: each dimension of gradient vector is derivative of output wrt value in that dimension
 - Transfer functions: Jacobians instead of derivatives
- Useful to know [matrix derivatives](#) notation

Transfer functions



$$J_f = \begin{pmatrix} \frac{\partial Y_1}{\partial X_1} & \frac{\partial Y_2}{\partial X_1} & \dots & \frac{\partial Y_n}{\partial X_1} \\ \frac{\partial Y_1}{\partial X_2} & \frac{\partial Y_2}{\partial X_2} & \dots & \frac{\partial Y_n}{\partial X_2} \\ \dots & \dots & \dots & \dots \\ \frac{\partial Y_1}{\partial X_m} & \frac{\partial Y_2}{\partial X_m} & \dots & \frac{\partial Y_n}{\partial X_m} \end{pmatrix} \quad \underline{v}_{out} = \begin{pmatrix} \frac{\partial Output}{\partial Y_1} \\ \frac{\partial Output}{\partial Y_2} \\ \dots \\ \frac{\partial Output}{\partial Y_n} \end{pmatrix}$$

- General function f

- $\underline{v}_{in} = J_f * \underline{v}_{out}$

- Linear function W

- $J_f = W^T$

- $\underline{v}_{in} = W^T * \underline{v}_{out}$

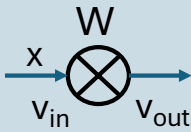
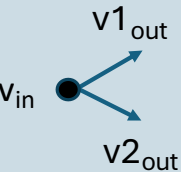
- Derivatives of output wrt weights in W (will be a matrix)

$$\frac{\partial(Output)}{\partial W} = \underline{v}_{out} \otimes \underline{x} \quad (\otimes \text{ is outer-product})$$

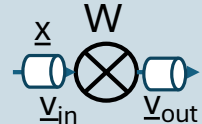
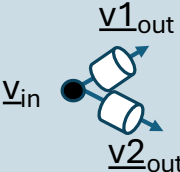
$$\frac{\partial(Output)}{\partial W}(i, j) = \underline{v}_{out}(i) * \underline{x}(j) \quad (\text{If } w(i, j) \text{ changes a small amount, how much does the output change?})$$

Transfer Functions (scalar and vector)

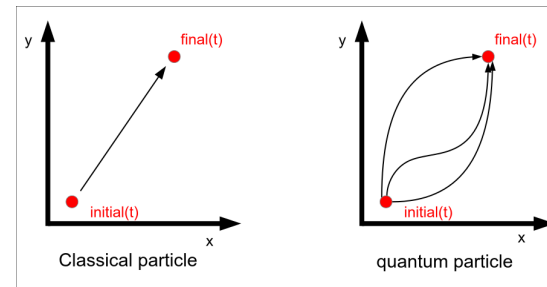
Scalar case

 $v_{in} = 1$	 $v_{in} = W * v_{out}$ $\frac{\partial (Output)}{\partial W_i} = v_{out} * x$
 $v_{in} = (v1_{out} + v2_{out})$	 $v1_{in} = \frac{\partial f}{\partial x}(x, y) * v_{out}$

Vector case

 $v_{in} = 1$	 $\underline{v}_{in} = W^T * \underline{v}_{out}$ $\frac{\partial (Output)}{\partial W} = \underline{v}_{out} \otimes \underline{x}$
 $\underline{v}_{in} = (\underline{v1}_{out} + \underline{v2}_{out})$	 $\underline{v}_{in} = J_f * \underline{v}_{out}$

Summary



Path integral formulation
of quantum mechanics (Feynman)

Abstraction for neural networks: parameterized programs

Gradient computation

- Abstractly (what?): sum over paths (important idea in Physics)
- Efficient computation (how?): compositional algorithm on dataflow graph representation

Handles complex nonlinear functions like $\sin(x)$, $\sin(\cos(x))$, $e^{\cos(x)}$

Handles complex neural networks with weight-sharing and irregular interconnections (such as “skip connections”) smoothly

With small caveat, handles conditional and loops as well

Example of backward dataflow analysis used in compilers

