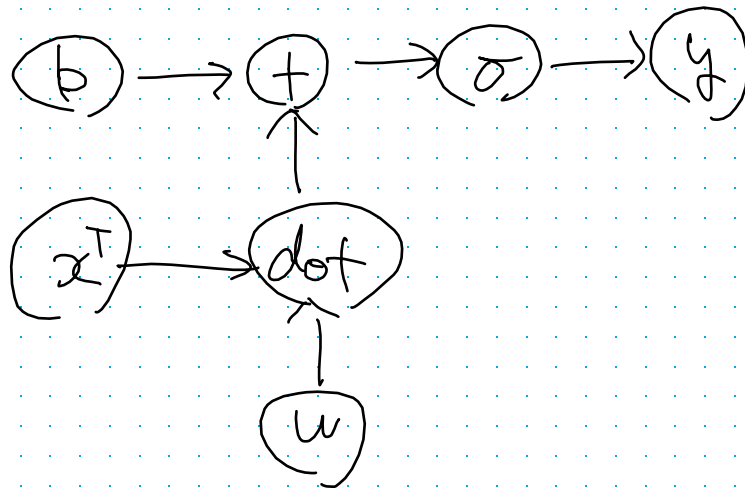


April 12th

$$(a) \hat{y} = \sigma(x^T w + b)$$

computation graph



(b) What kind of graph is computational graph in general?

→ DAG (Directed Acyclic Graph)

except for RNN.

How many edges can it have at most?

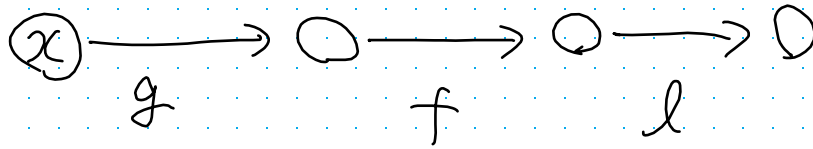
Assume N node

$$(N-1 + N-2 + \dots + 0) = (N-1) \times \frac{N}{2}$$



April 12

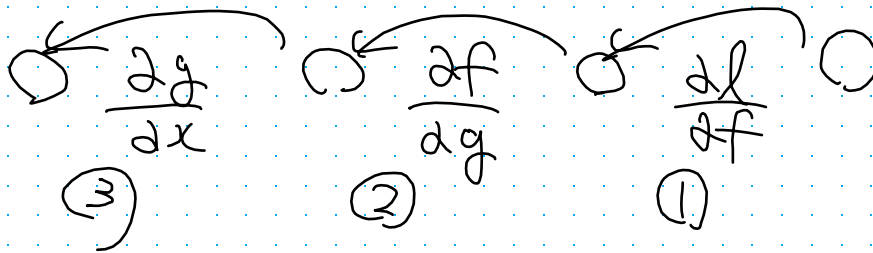
(c) Explain how the back prop algo find the grad



$$\nabla(l \circ f \circ g)(x)$$

$\frac{\partial l}{\partial f} \frac{\partial f}{\partial g} \frac{\partial g}{\partial x}$ can be computed $x \in \mathbb{R}^n$ as follows

Back propagation



④ using chain rule

$$\nabla(l \circ f \circ g)(x) = \frac{\partial l}{\partial f} \frac{\partial f}{\partial g} \frac{\partial g}{\partial x}$$

⑤ Loop num of output dimension

in this case
output $\in \mathbb{R}$
then
we don't
have to
accumulate.

(d) What is the order of computation required to compute the grad

let T as in this case
number of transition, in other words
edge.

let D as
dimension of output

Then Order of computation of
backward is

$$\underline{O(TD)}$$

April 12th

$$(e) \quad l: \mathbb{R}^n \rightarrow \mathbb{R}^k \quad k \gg 1$$

which mode is prefer to compute $\mathcal{D}(\dots)$
efficiently

Answer

forward mode automatic diff
is preferable

if $k > n$

otherwise backward mode
automatic differentiation
is preferable

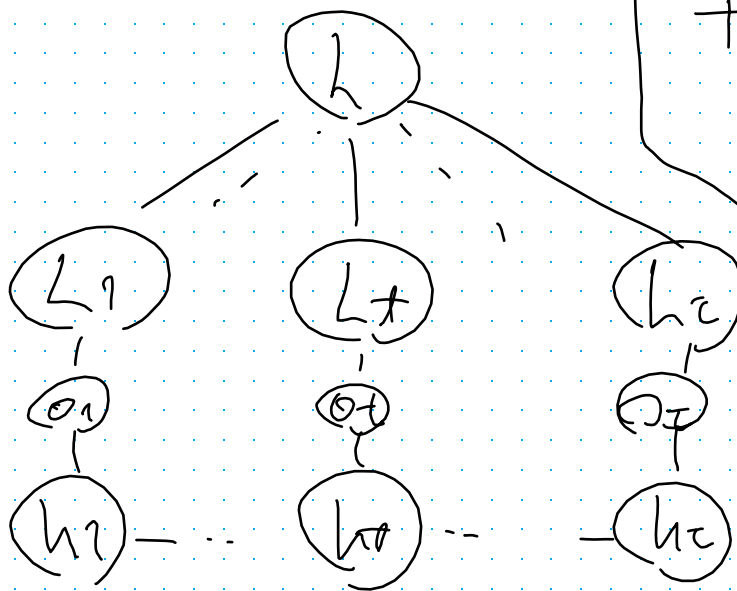
because forward mode AD
computational cost of

depends on dim of input (n)

Similarly, reverse mode AD's cost
depends on dim of output (k).

(f) (Bonus) Which node of AD can you use to train an RNN without storing the activation?

How does it work?



The drawback of this approach is computational cost

depends on dim of input
in general RNN has large dim of output.

Without storing activation

we have to use forward node AD.

Which can be computed

at the same time of forward process and doesn't require to store activation

(g) Bonus

$$x_{t+1} = x_t - H^{-1} \nabla \ell(x_t)$$

How can you combine the two mode of AD to compute Hv efficiently for any vector v ?

Ans Hessian vector product doesn't calculate H explicitly

step 1. compute grad and store activation as reverse mode at t step

step 2 as a forward process using input as v_{t+1} and compute Hv by using t step grad.

April 12

(from JAX Tutorial)

$$f: \mathbb{R}^n \rightarrow \mathbb{R}$$

The Hessian at a point $x \in \mathbb{R}^n$ is written as $d^2f(x)$.

A Hessian-vector-product function is then able to evaluate

$$u \mapsto d^2f(x) \cdot u \quad \text{for any } u \in \mathbb{R}^n$$

$$d^2f(x) u = d [x \mapsto df(x) \cdot u] = dg(x)$$

where $g(x) = df(x) \cdot u$

$$\text{grad} \left(\underbrace{\text{jnp.dot}}_{\text{jax numpy}} (\text{grad}(f(x)), u) \right)$$

$$(x, u) \mapsto \partial^2 f(x) u = \partial g(x) \cdot u$$

$$g: \mathbb{R}^n \rightarrow \mathbb{R}^n, \quad g(x) = \partial f(x) \quad \rightarrow \quad \frac{\partial}{\partial x} \left(\underbrace{\partial f(x)}_{(2)} \cdot u \right)_{(3)}$$

$$\circ \text{ findrev_hvp}(f, x, u)$$

$$\text{jvp}(\partial f, x, u)$$

$$\circ \text{ revfwd_hvp}(f, x, u)$$

$$g = \text{jvp}(f, x, u)$$

$$\text{grad}(g, x)$$

find τ
 $\partial^2 f(x) u = \tau$
 $1 = \tau^T \tau < \tau^T \tau$
 交叉的

$$\rightarrow \frac{\partial}{\partial x} \left(\underbrace{\partial f(x)}_{(2)} \right) \cdot u_{(3)}$$

JVP

JVP (Jacobian-Vector Products)

$$f: \mathbb{R}^n \rightarrow \mathbb{R}^m \quad \partial f(x) \in \mathbb{R}^{m \times n}$$

$$(x, u) \mapsto \partial f(x) u$$

~~Efficiently~~

$\nabla f(x) \cdot u \leftarrow$ compute by forward mode
 setting $\dot{x} = u$

$\nabla^2 f(x) \cdot u \leftarrow$ take the result and
 apply backprop

[Pros]

it allows to evaluate a function
 and its derivative at the same time.

$\Downarrow$

This enable to overcome the two phase doctrine
 of reverse mode. (救済)

$\uparrow$

This comes at cost of storing
 dual representation of all intermediate
 node.

<reverse-on-forward> Hyp

if 1st order 計算とあわせて
 forward 2nd order 計算と連ねる!

HVP

< key insight >

defining $g(w) = \nabla J(w) = \frac{\partial f}{\partial w}$

then H is just the Jacobian of g

⇓

This leads to an HVP implementation
called forward-over-reverse
AD

How to take Hvp $\nabla^2 f(x) \cdot u$

Notation $jvp(f, x, u)$

$$\left(= Dg \cdot u - \frac{d}{dt} Dg(x + tu) \right)?$$

$$\boxed{(x, u) \mapsto \nabla f(x) \cdot u} \leftarrow \text{JAX} \quad \text{2nd Definition}$$

Reverse - over - Reverse

$$\nabla(x \mapsto \nabla f(x) \cdot u)$$

Forward - over - Reverse

$$\frac{d}{dt} D(\nabla f)(x + tu) = jvp(\nabla f, x, u)$$

Reverse - over - Forward

$$\nabla \left[\frac{d}{dt} (Df)(x + tu) \right] = \nabla jvp(f, x, u)$$

$$\text{Jvp}(f, x, u): (x, u) \mapsto df(x) \cdot u$$

$$\text{grad}(f, x) \quad x \mapsto \nabla f(x)$$

1. reverse-over-reverse

$$[\text{computation}] \rightarrow \text{grad}(\text{dot}(\text{grad}(f, x), u), x)$$

$$\nabla_x(\nabla_x f(x) \cdot u) = \nabla_x^2 f(x) \cdot u$$

2. forward-over-reverse

$$[\text{computation}] \rightarrow \text{Jvp}(\text{grad}(f, x), x, u)$$

$$\partial \nabla_x f(x) \cdot u = \nabla_x^2 f(x) \cdot u$$

3. reverse-over-forward

$$[\text{computation}] \rightarrow \text{grad}(\text{jvp}(f, x, u), x)$$

$$\nabla_x(df(x) \cdot u) = \nabla_x^2 f(x) \cdot u$$