

Deep Learning

Grundlagen Grundbegriffe Grundübungen

Alfatraining Kurs 2025-08 Dozent: Karsten Flügge

Themen des Tages

Tag 12

Recurrent Neural Networks (RNN) Teil 2:
Exploding und Vanishing Gradient Probleme #2
LSTM (Long Short-Term Memory)
GRU (Gated Recurrent Unit)
Deep RNN, Deep LSTM

Recurrent Neural Networks (RNN) II

nächster output vom RNN:

$$\mathbf{h}_{t+1} = \sigma(\mathbf{W} \cdot \mathbf{x} + \mathbf{W}_h \cdot \mathbf{h}_t + \mathbf{b})$$

$\mathbf{W}$ Gewichtsmatrix für Input $\mathbf{x}$

$\mathbf{W}_h$ Gewichtsmatrix für Inneren Zustand 'hidden' $\mathbf{h}_t$

$\mathbf{b}$ bias

=> Mathematik kann übersichtlicher sein als Code

Recurrent Neural Networks (RNN) II

Achtung: "number of layers" $\neq$ sequence length

$h_{t+1} = \sigma(W \cdot x + W_h \cdot h_t + b) \Rightarrow$ into next layer:

$h2_{t+1} = \sigma(W2 \cdot h1_t + W2_h \cdot h2_t + b2)$

$h3_{t+1} = \sigma(W3 \cdot h2_t + W3_h \cdot h3_t + b3)$

Die RNN Schichten werden also nicht nur vertikal sondern auch horizontal gestapelt. Die Ausgabe ist dann das was rechts oben herauskommt.

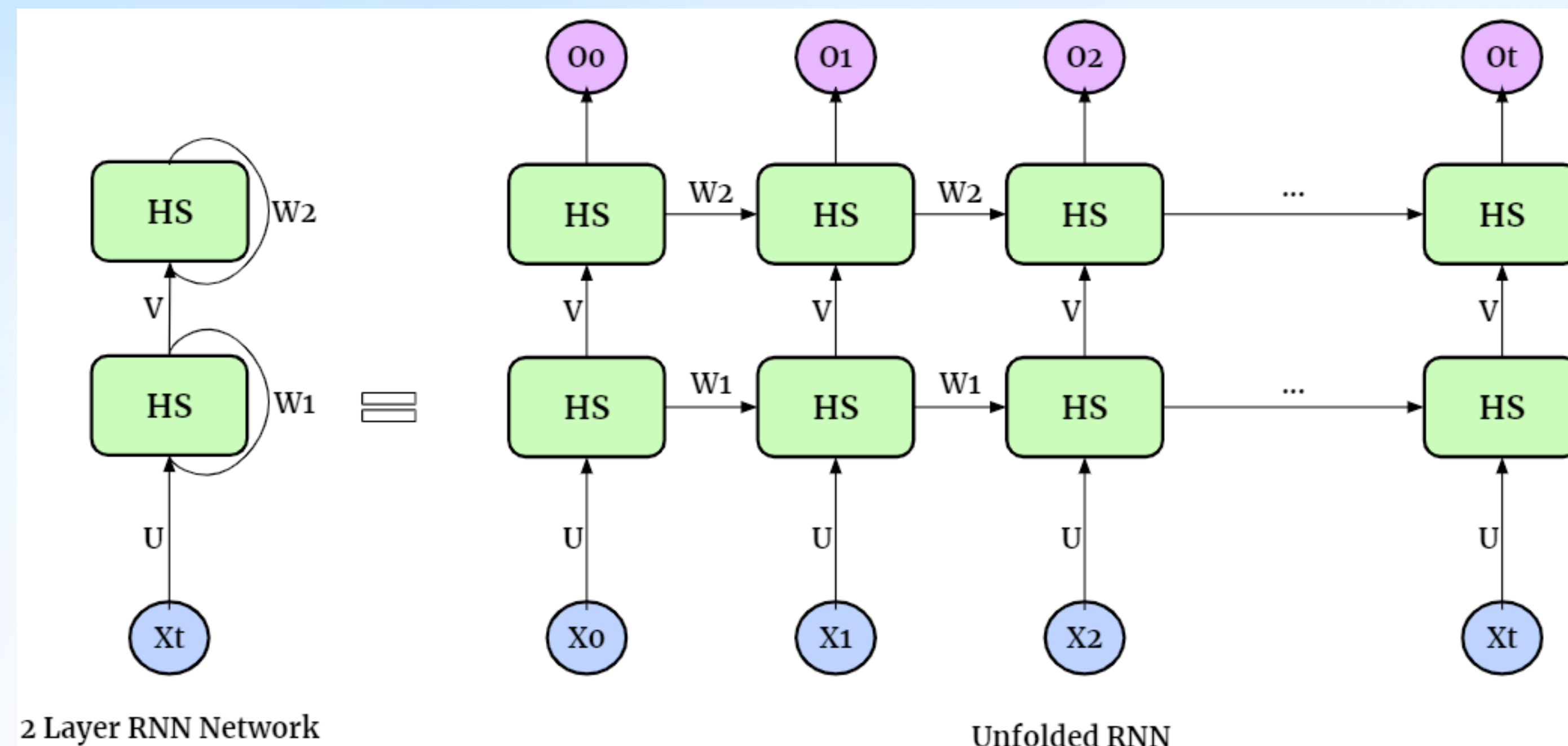
Recurrent Neural Networks (RNN) II

Achtung: "number of layers" $\neq$ sequence length

$h_{t+1} = \sigma(W \cdot x + W_h \cdot h_t + b) \Rightarrow$ into next layer:

$h2_{t+1} = \sigma(W2 \cdot h1_t + W2_h \cdot h2_t + b2)$ # memory 2

$h3_{t+1} = \sigma(W3 \cdot h2_t + W3_h \cdot h3_t + b3)$



Exploding und Vanishing Gradients in RNN

Wir erinnern uns an Exploding und Vanishing Gradient Probleme. Diese treten bei RNNs besonders auf. Lösung:

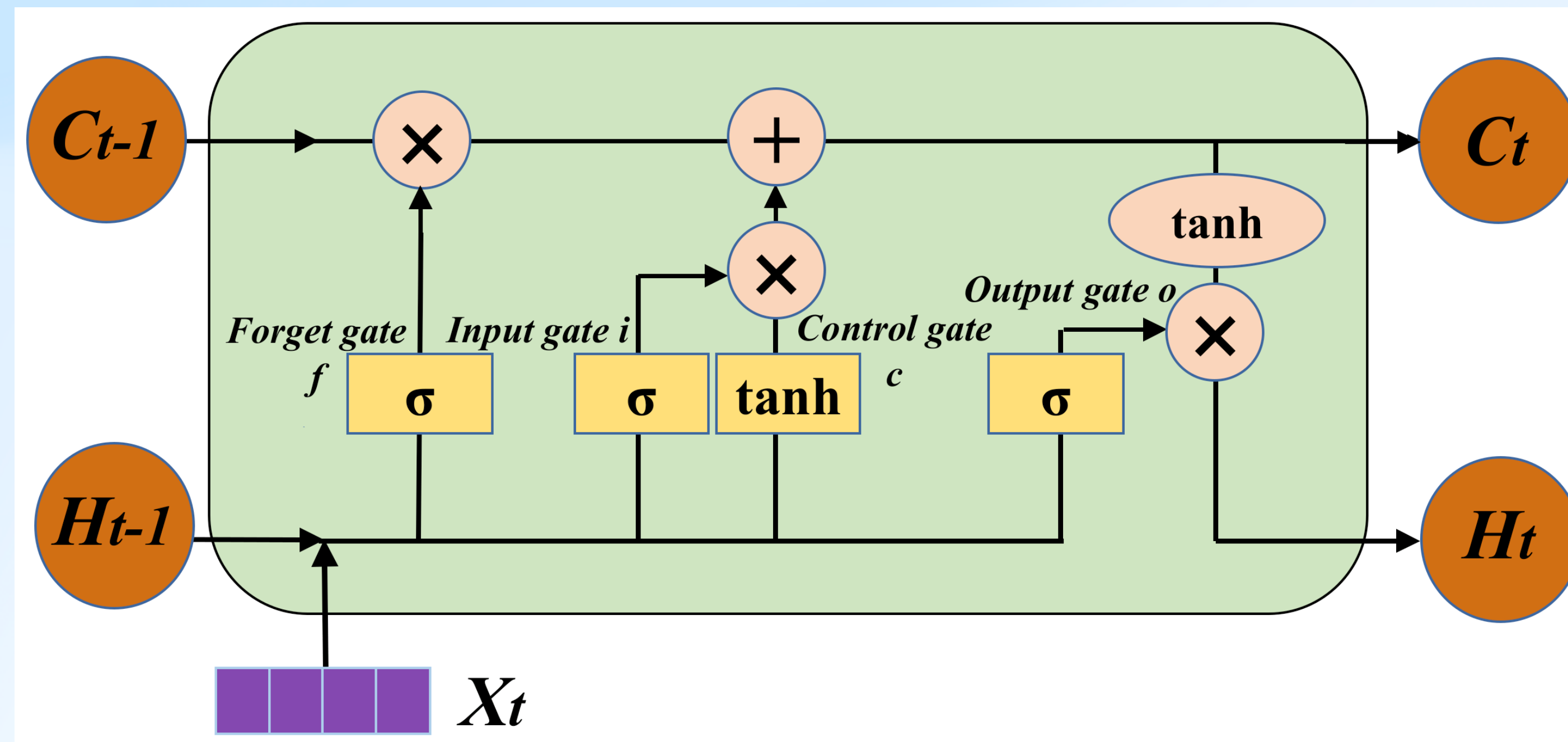
Modifizierte RNNs:

LSTM (Long Short-Term Memory)

Ein LSTM besteht aus sogenannten LSTM-Zellen, die über die Zeit hinweg verbunden sind.

Jede **LSTM-Zelle** hat drei Hauptkomponenten:

- **Vergessen-Gate (Forget Gate)**: Bestimmt, welche **Informationen** aus der Zelle **entfernt** werden sollen.
- **Eingangs-Gate (Input Gate)**: Entscheidet, welche neuen **Informationen hinzugefügt** werden.
- **Ausgangs-Gate (Output Gate)**: Bestimmt, welche **Informationen** aus der Zelle **weitergegeben** werden.

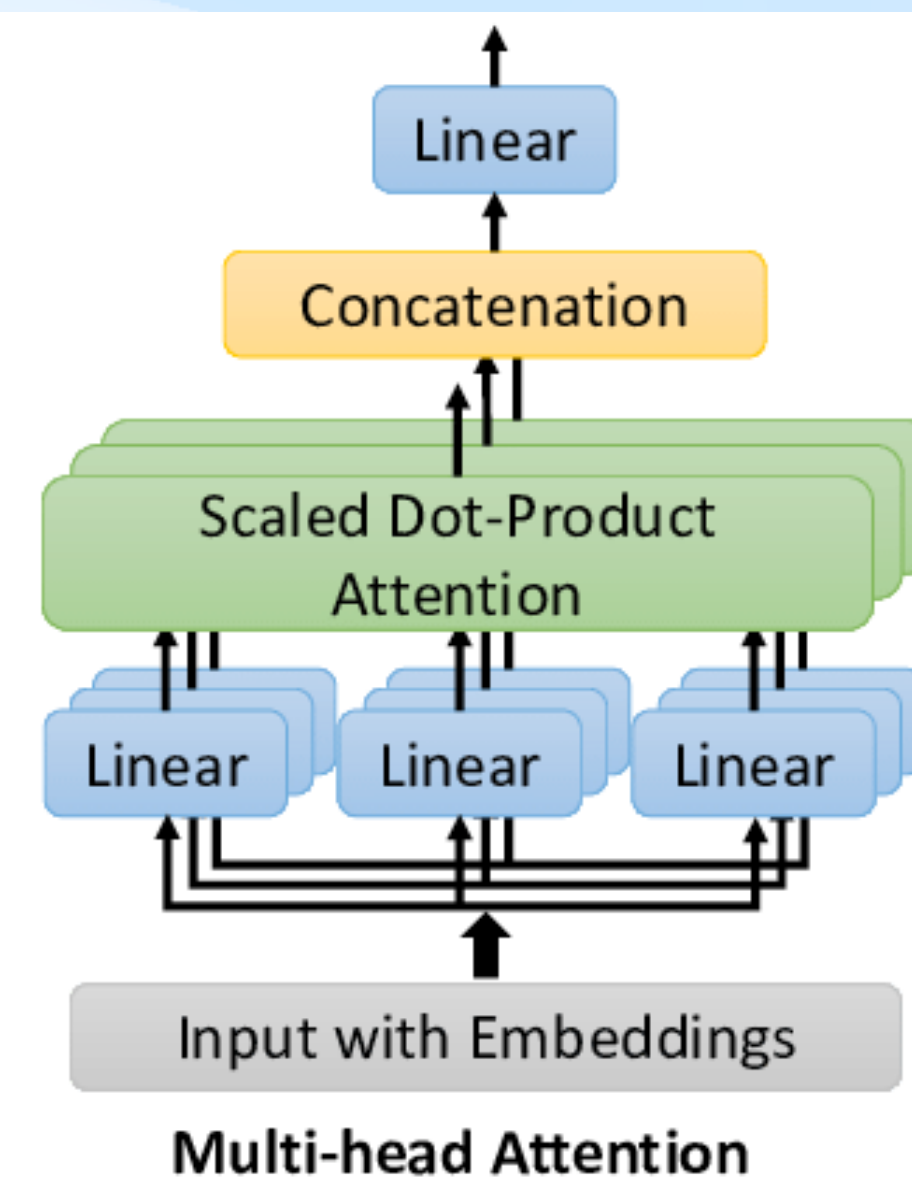
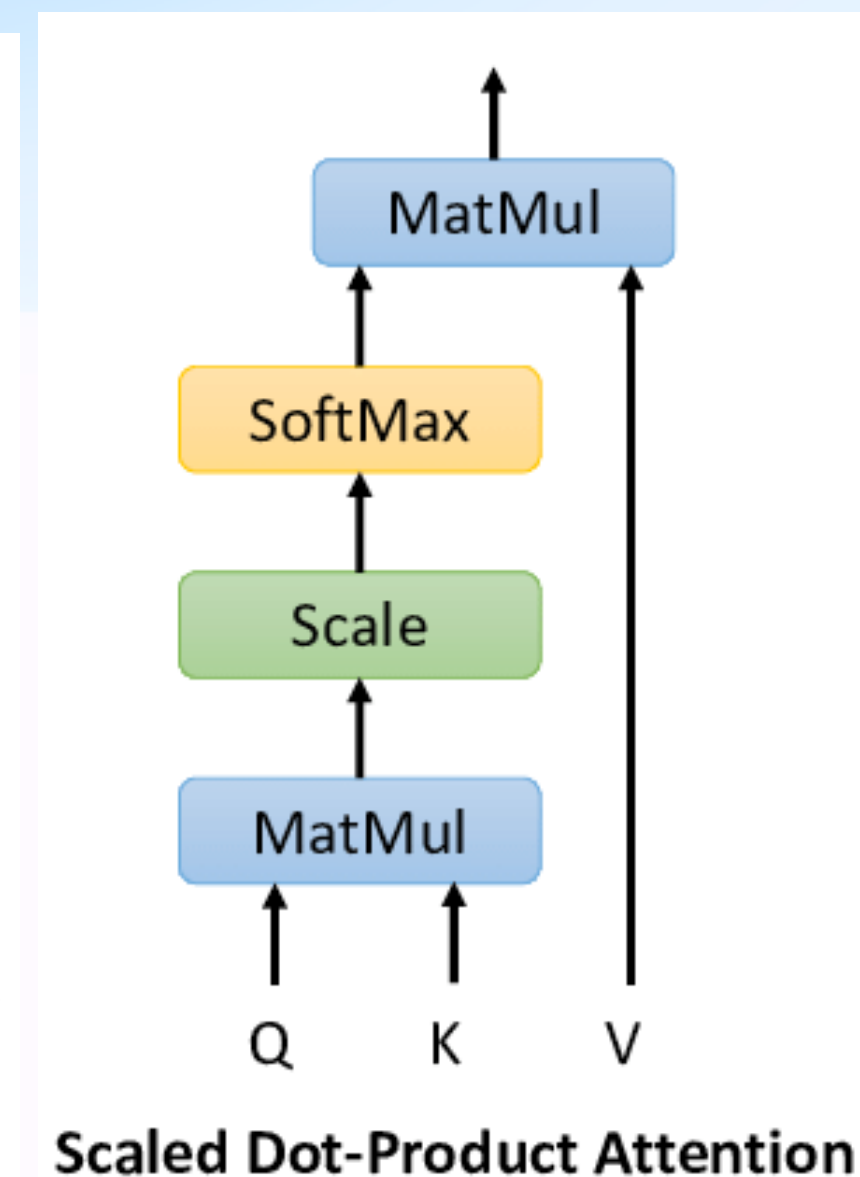
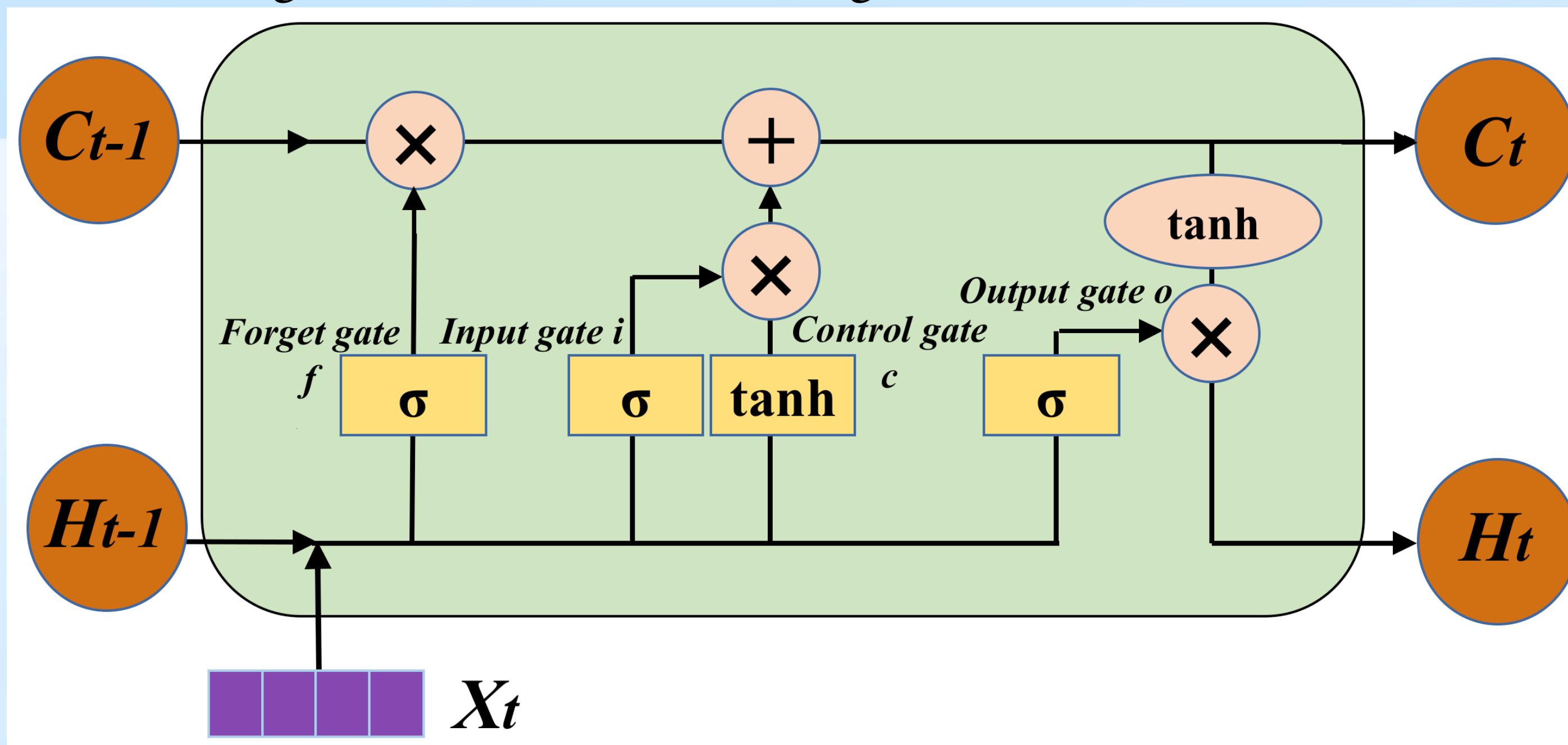


LSTM vs Attention

- Entsprechen die Berechnungen im LSTM UNGEFÄHR den Attention?
- **Vergessen-Gate (Forget Gate)**: Bestimmt, welche **Informationen** aus der Zelle **entfernt** werden sollen.
- **Eingangs-Gate (Input Gate)**: Entscheidet, welche neuen **Informationen hinzugefügt** werden.
- **Ausgangs-Gate (Output Gate)**: Bestimmt, welche **Informationen** aus der Zelle **weitergegeben** werden.

LSTM: sigmoid 'conditions/gates'

Attention: sigmoid + softmax 'condition/gate'



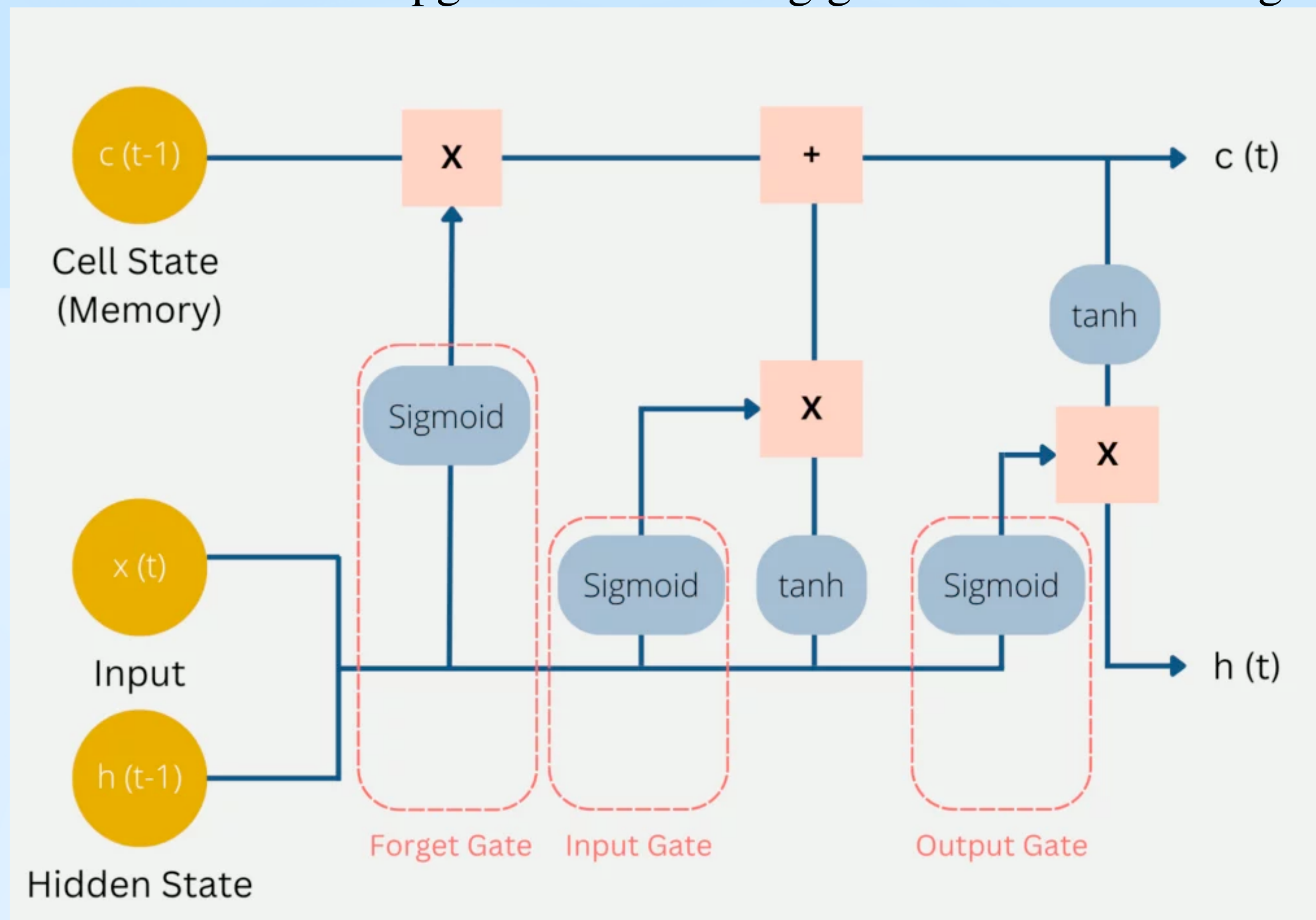
LSTM (Long Short-Term Memory)

Anders als bei dem einfachen RNN werden beim LSTM **zwei zeitabhängige Zustände** gespeichert:

Cell State entspricht so etwas wie **Langzeitgedächtnis**

Hidden State entspricht etwa **Kurzzeitgedächtnis** und der **aktuellen Aktivierung** (⚠ anders als normale Konvention)

In jedem Zeitschritt werden die Werte upgedated in Abhängigkeit von den vorherigen Werten und dem aktuellen Input:



LSTM (Long Short-Term Memory)

Idee als Rechenschema:

für jeden Zeitschritt t :

Forget-Gate $f_t = \sigma(W_f * [h_{t-1}, x_t] + b_f)$

Input-Gate $i_t = \sigma(W_i * [h_{t-1}, x_t] + b_i)$

Candidate-Zustand $C'_t = \tanh(W_c * [h_{t-1}, x_t] + b_c)$

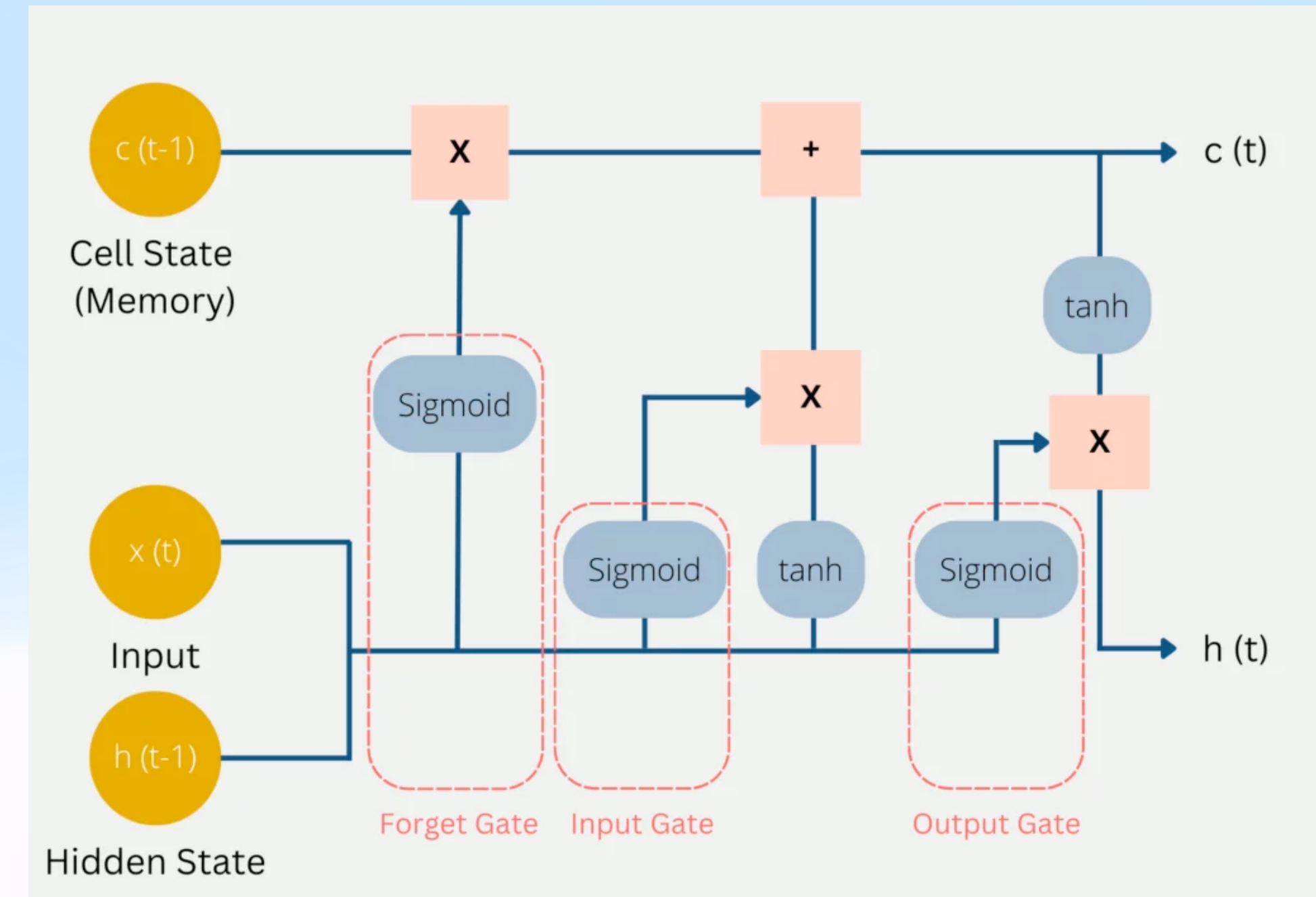
Cellzustand $C_t = f_t * C_{t-1} + i_t * C'_t$

Output-Gate $o_t = \sigma(W_o * [h_{t-1}, x_t] + b_o)$

Cell-Ausgabe $h_t = o_t * \tanh(C_t)$

zum Vergleich einfacheres RNN Schema:

$$h_{t+1} = \sigma(W \cdot x + W_h \cdot h_t + b)$$



Deep RNN, Deep LSTM

Deep RNN, Deep LSTM

wie Lego-Blöcke lassen sich RNN **Schichten** / LSTM **Zellen** hintereinander schalten, indem man schlicht die Ausgabe einer Zelle als Eingabe für die nächste Zelle verwendet.

Natürlich darf man die Ausgaben zwischendurch auch mit anderen Netzwerkkomponenten bearbeiten, besonders linear / fully connected layers. Wir haben auch gesehen wie man die Ausgaben splitten und später wieder re-kombinieren kann. Der Kreativität sind quasi keine Grenzen gesetzt? Doch: normale RNNs leiden unter Vanishing Gradient Problemen, Aber: dank LSTM und anderen Techniken welche wir gesehen haben lassen sich diese Grenzen doch meist überwinden.

Tatsächlich gab es und gibt es Ansätze, statt sich ad-hoc Netzwerk Architekturen auszudenken, diese lieber systematisch von Computern selber erforschen zu lassen.

Deep RNN, Deep LSTM

- ResNet: Often used with depths of 50, 101, or 152 layers.
- DenseNet: Commonly used with 121, 169, or 201 layers.
- EfficientNet: Variants range from 25 to 66 layers.
- Transformers (e.g., BERT): Often have depths of 12 to 48 layers or more.
 - **GPT-2**: 48 layers.
 - **GPT-3**: 96 layers.
 - **GPT-4**: wahrscheinlich 100te Schichten (Details proprietär)

LSTM already outdated?

Apropos Kreativität:

LSTM Rechenschema ist **kompliziert**, was ist tatsächlich die **Essenz** dieser Zellen?

Memory, in Abhängigkeit vom vorherigen Zustand und vom aktuellen Input

Warum darf man diesen Zustand nicht einfach irgendwo anders im Speicher oder auf der Festplatte abspeichern?

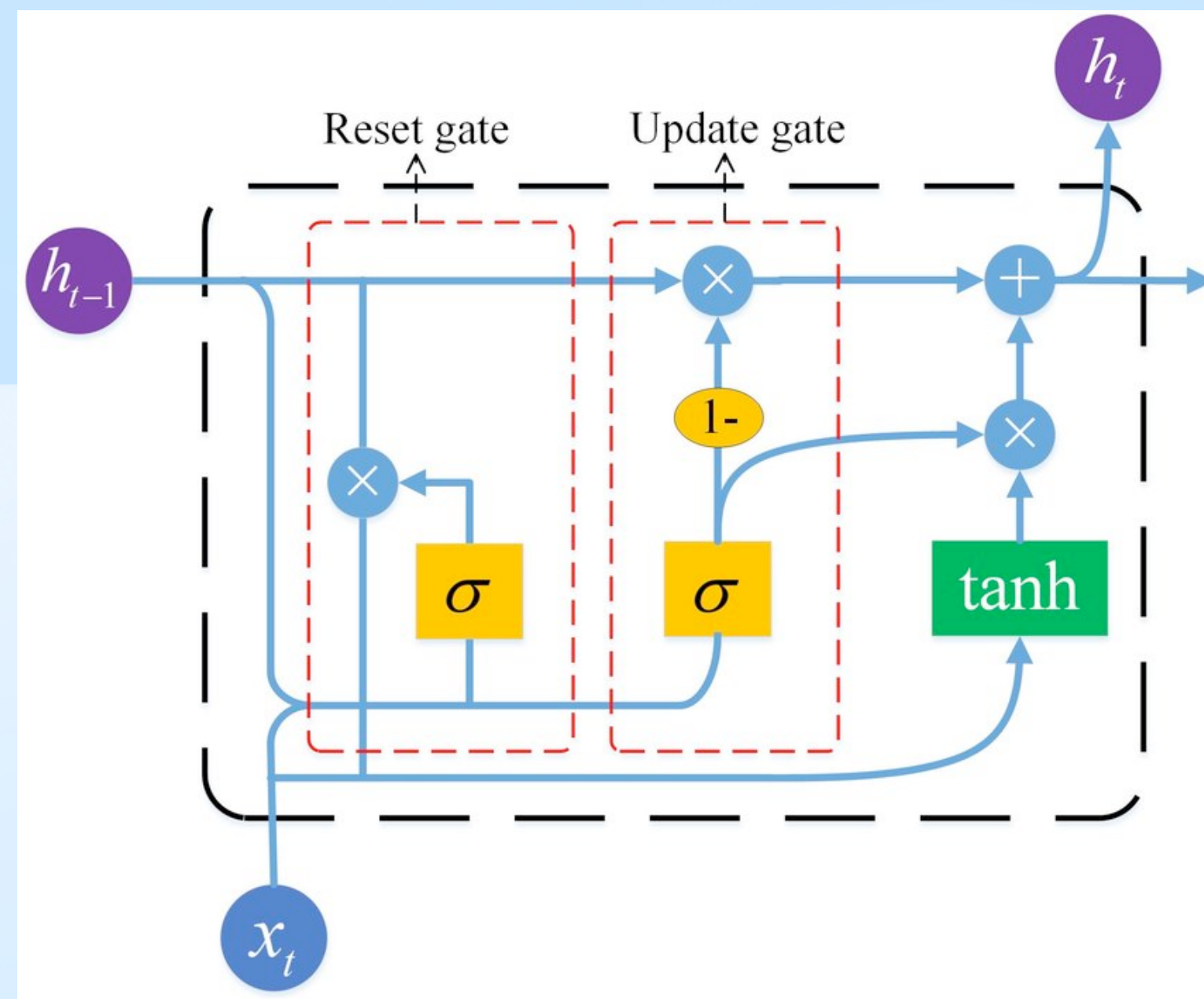
Um trainierbar zu sein muss jede Aktion in einem neuronalen Netz ableitbar sein oder zumindest brauchen wir so etwas für einen Pseudo **Gradienten**. Es gibt Forschungsansätze solche **allgemeinen Schreib und Lesevorgänge** als ableitbare Funktion zu approximieren, zunächst gucken wir aber auf eine Vereinfachung des Rechenschemas:

GRU (Gated Recurrent Unit)

GRU (Gated Recurrent Unit): Vereinfachung von LSM und Beibehaltung der selben Grundfunktion:
Zumindest für tiefe Netze hat sich herausgestellt dass GRUs keinerlei Genauigkeitsverluste gegenüber LSTMs haben.

nur zwei Gates

- **Reset Gate:** Entscheidet, wie stark der vorherige Zellzustand **vergessen** wird.
- **Update Gate:** Kontrolliert den Anteil der alten Informationen, die **beibehalten** werden, und den Anteil der neuen Informationen, die **hinzugefügt** werden.



Im Grunde sind wir also fast zurück beim RNN, an der Ausgabe Schnittstelle ist es sogar vereinfacht weil Hidden State und Aktivierung zusammenfallen.
Nur intern sind GRUs etwas inspirierter als klassische RNNs:

GRU (Gated Recurrent Unit)

Für jeden Zeitschritt t :

Reset-Gate:

$$\mathbf{r}_t = \sigma(\mathbf{W}_r \cdot [\mathbf{h}_{t-1}, \mathbf{x}_t] + \mathbf{b}_r)$$

$\mathbf{W}_r$: Gewichtsmatrix für das Reset-Gate

$\mathbf{b}_r$: Bias für das Reset-Gate

$\mathbf{h}_{t-1}$: Hidden-State vom vorherigen Zeitschritt

$\mathbf{x}_t$: Eingabe im aktuellen Zeitschritt

$[\mathbf{h}_{t-1}, \mathbf{x}_t]$ concat von hidden und input

σ : **Sigmoid**-Funktion

Update-Gate:

$$\mathbf{z}_t = \sigma(\mathbf{W}_z \cdot [\mathbf{h}_{t-1}, \mathbf{x}_t] + \mathbf{b}_z)$$

selbe Formel nur andere Verwendung

$\mathbf{W}_z$: Gewichtsmatrix für das Update-Gate

$\mathbf{b}_z$: Bias für das Update-Gate

Kandidaten-Aktivierung "new hidden"

$$\tilde{\mathbf{h}}_t = \tanh(\mathbf{W}_h \cdot [\mathbf{r}_t * \mathbf{h}_{t-1}, \mathbf{x}_t] + \mathbf{b}_h)$$

$\mathbf{W}_h$: Gewichtsmatrix für die Kandidaten-Aktivierung

$\mathbf{b}_h$: Bias für die Kandidaten-Aktivierung

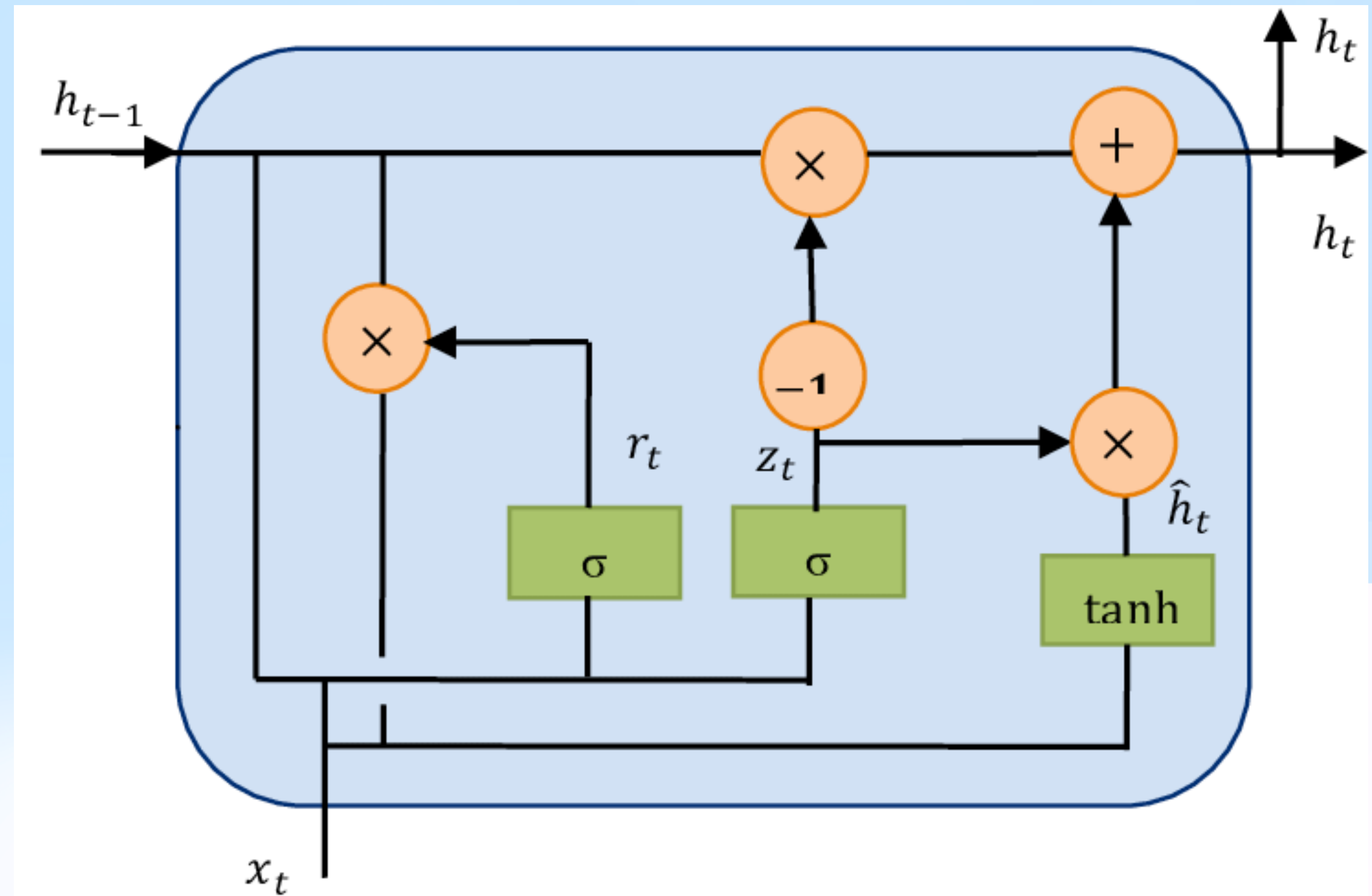
$\mathbf{r}_t * \mathbf{h}_{t-1}$: Elementweise Multiplikation des Reset-Gates mit dem Hidden-State vom vorherigen Zeitschritt ($\mathbf{r}=0 \Rightarrow$ old hidden verschwindet $\mathbf{r}=1$ keep it)

Aktualisierter Hidden-State:

$$\mathbf{h}_t = (1 - \mathbf{z}_t) * \mathbf{h}_{t-1} + \mathbf{z}_t * \tilde{\mathbf{h}}_t \quad \text{in } [-1, 1]$$

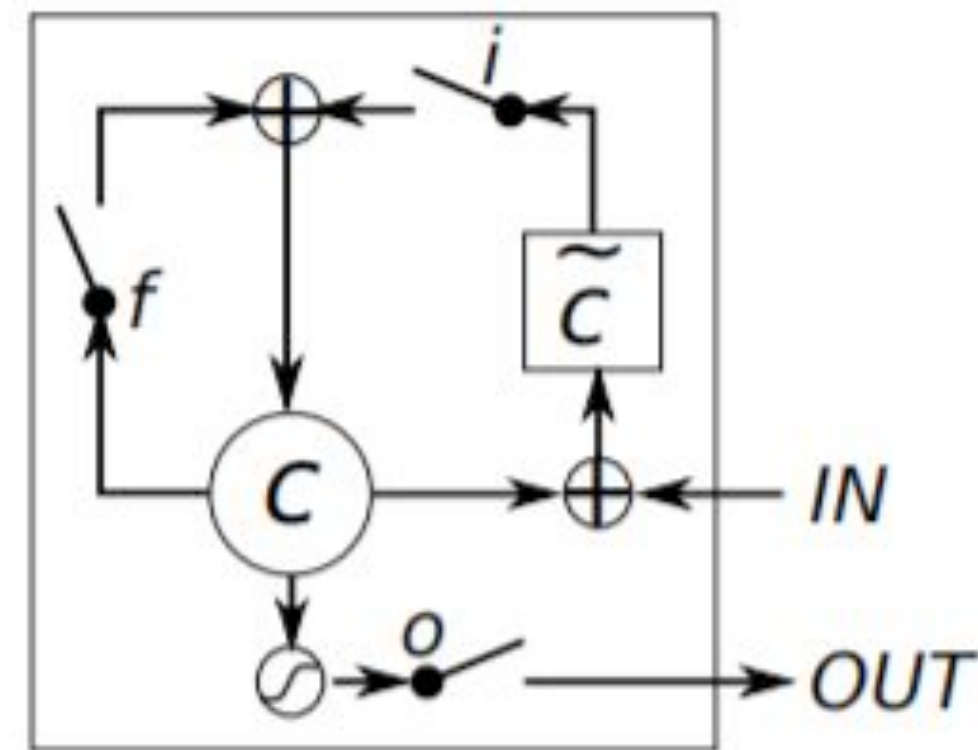
$(1 - \mathbf{z}_t) * \mathbf{h}_{t-1}$: Beibehaltung eines Teils des vorherigen Hidden-States

$\mathbf{z}_t * \tilde{\mathbf{h}}_t$: Hinzufügen der **neuen** Informationen aus der Kandidaten-Aktivierung

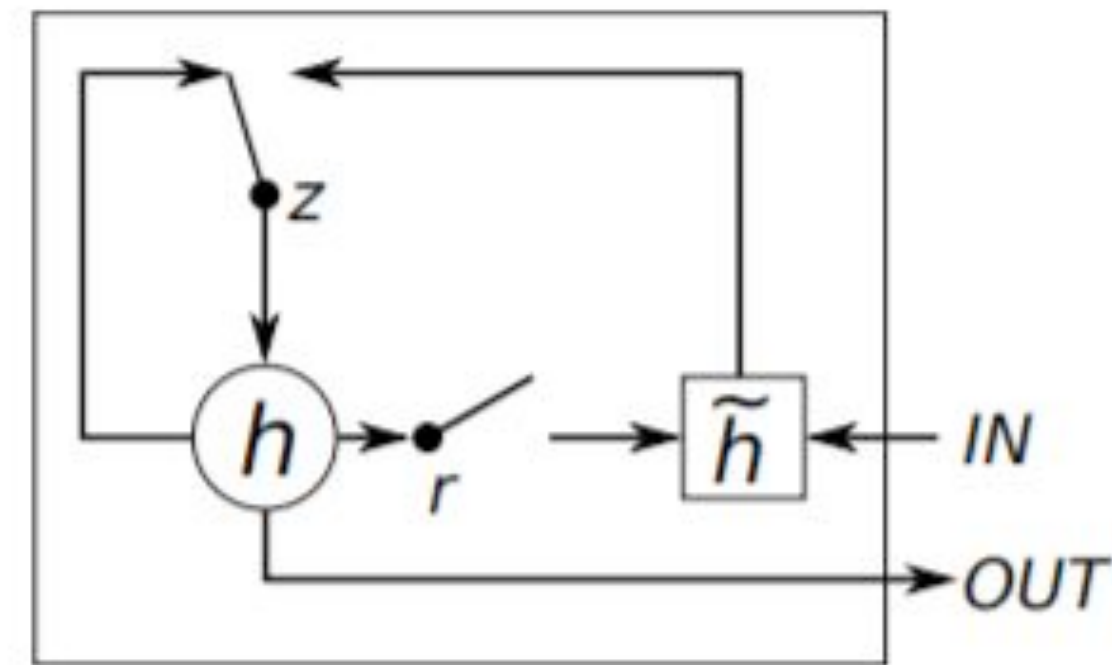


GRU (Gated Recurrent Unit)

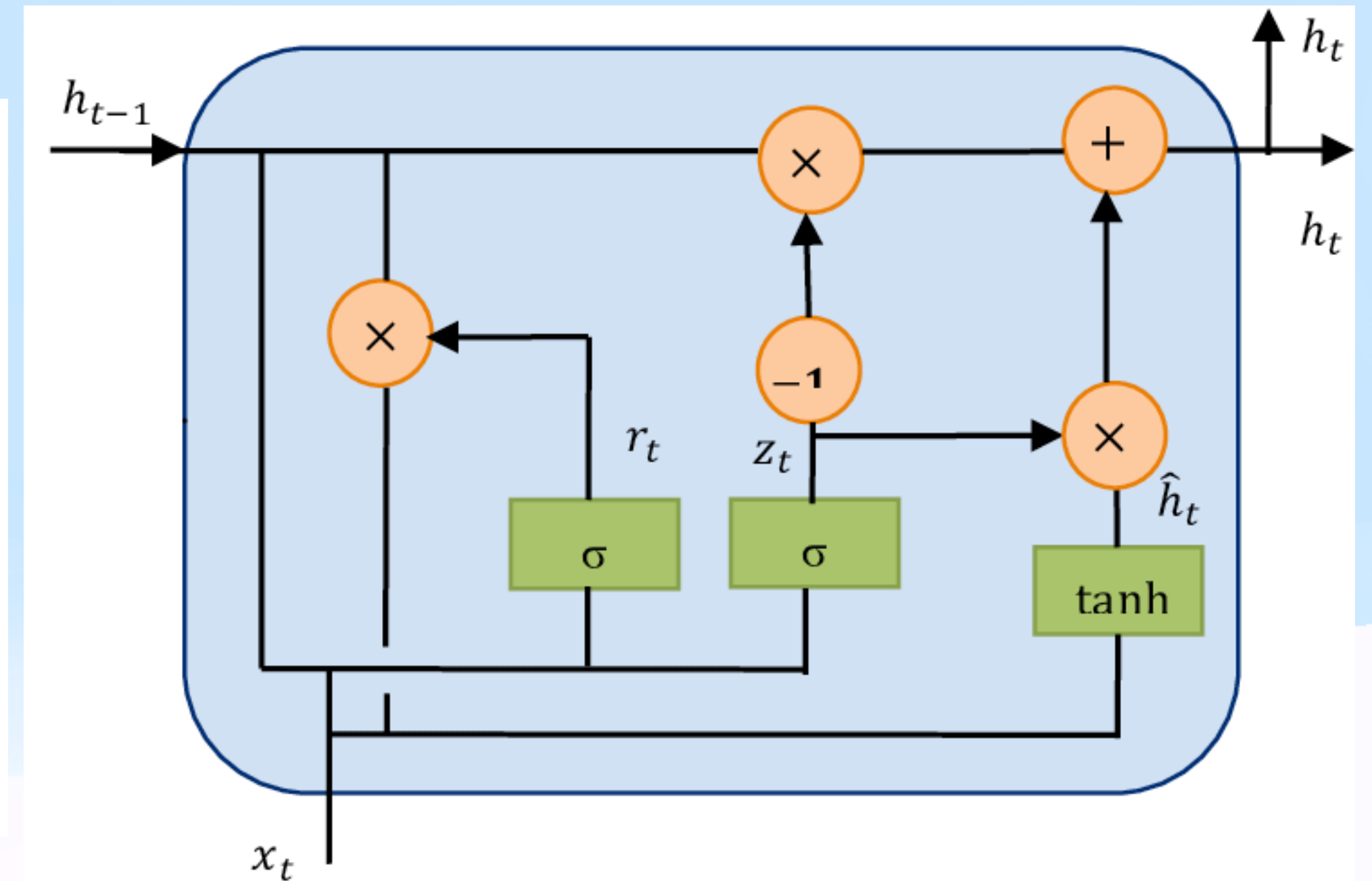
GRU gates $\approx$ logic gates



(a) Long Short-Term Memory



(b) Gated Recurrent Unit



Shakespeare RNN

Shakespeare RNN: Lasse Netzwerk schreiben wie Shakespeare

One hot encoding für Text: ASCII Buchstabe 'a' $\Rightarrow$ 97 'D' $\Rightarrow$ 68 ...

Jeder Buchstabe bekommt eine 1 im Vector:

$(0,0,0,0,0,0,0,0,0, \dots 1, 0,0,0,0,0,0, \dots 0) = \text{'a'}$ wenn 1 an 97ster Stelle

Buchstabe zu Vector

Text zu Tensor: Stapel von Buchstaben Vektoren

Damit kann unser Netz mit Texten rechnen.

Output Vorhersage = nächster Buchstabe

Zeitreihen langfristige Abhängigkeiten

QUEEN ELIZABET_ sage 'H' voraus aufgrund *aller* vorheriger Zeichen

Char RNNs

The Unreasonable Effectiveness of Recurrent Neural Networks

Andrej Karpathy blog May 21, 2015

<http://karpathy.github.io/2015/05/21/rnn-effectiveness/>

pure python implementation: **char_rnn_pure.py** !

pytorch implementation: **char_rnn.py**

<https://github.com/spro/char-rnn.pytorch>

```
self.encoder = nn.Embedding(input_size, hidden_size)
self.rnn = nn.GRU(hidden_size, hidden_size, n_layers) # << das ist alles!
self.decoder = nn.Linear(hidden_size, output_size)
```

Char RNNs

The original RNN produces 'somewhat english' output after a million steps:

DUKE VINCENTIO:

Maday; that the present for whitted it is fair have up my ours I am you and thy broose,
Bans.

Lord.

Warkerd him on

Your queen royall a goods to oned your field

Modern variants with GRU produce 'almost Shakespeare' results after only a few iterations:

Why: my lord, unto them like

By her a loves to the sky.

PROSPERO:

A foul mother's liege and all work