

Deep Learning

Grundlagen Grundbegriffe Grundübungen

Alfatraining Kurs 2025-08 Dozent: Karsten Flügge

Themen des Tages

Tag 13

Textverarbeitung durch neuronale Netze

Text-Preprocessing

Embedding-Schichten

Text-Klassifizierung

Sentiment-Analyse

Transfer Learning in NLP

Textverarbeitung durch neuronale Netze

Die **Textverarbeitung** durch neuronale Netze, auch bekannt als **Natural Language Processing (NLP)**, ist ein zentrales Anwendungsgebiet des maschinellen Lernens, insbesondere tiefen neuronalen Netzen (Deep Learning). NLP befasst sich mit der Interaktion zwischen Computern und **menschlicher Sprache**. Es ermöglicht Maschinen das **Verständnis**, die **Interpretation** und die **Generierung** von natürlicher Sprache in einer Weise, die nützlich ist.

Textverarbeitung durch neuronale Netze

Anwendungsbereiche:

- **Maschinelle Übersetzung** (z.B. Google Translate)
- **Textklassifikation** (z.B. Spam-Erkennung, Sentiment-Analyse)
- **Spracherkennung** (z.B. Sprachassistenten wie Siri oder Alexa)
- **Voice preservation & Deep Fakes: voice cloning**
- **Textgenerierung** (z.B. GPT-Modelle wie dieses hier)
- **Zusammenfassung von Texten** (z.B. von Artikeln)
- **Frage-Antwort-Systeme** (z.B. Chatbots)

Text-Preprocessing

- Wie kann ein neuronales Netz mit Text 'rechnen'?

Naiver Ansatz: ASCII als one-hot-encoding input tensor mit 128/256 Breite
256 Zeichen => [0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 1 an stelle 32 == "Space", 0,0,0]

oder nur printable's :
'0123456789abcdefghijklmnopqrstuvwxyzABCDEFGHIJKLMNOPQRSTUVWXYZ!"#\$%&\'()*+,-./:;<=>?@[\\]^_`{|}~ \t\n\r\x0b\x0c'

Dec	Hex		Dec	Hex		Dec	Hex		Dec	Hex		Dec	Hex		Dec	Hex		Dec	Hex	
0	00	NUL	16	10	DLE	32	20		48	30	0	64	40	@	80	50	P	96	60</	

Zeichen-basierte Tokenisierung

- **Wie kann ein neuronales Netz mit Text 'rechnen'?**

Naiver Ansatz: ASCII/Bytes als one-hot-encoding input tensor mit 128/256 Breite

- **Nachteil:**
 - Ist recht **verschwenderisch** weil viele AscII Codes gar nicht vorkommen.
 - Größere Sequenzen, was die Modellgröße und Rechenzeit erhöht.
 - Die **UTF-8** Erweiterung macht einige Codes überladen (braucht Kontext)
- **Vorteile:**
 - sehr einfache Implementierung
 - Vermeidung von Artefakten anderer Methoden:
 - **GPT** kann schlecht Worte / Buchstaben zählen + rechnen
 - “Vorzeitige Optimierung ist die Wurzel allen Übels.”
 - Lernen von Daten jeglicher Art (Bilder in ChatGPT)

Preprocessing

Filtern von gewünschten Zeichen

```
import string
all_characters = string.printable
n_characters = len(all_characters) # 100
```

```
'0123456789abcdefghijklmnopqrstuvwxyzABCDEFGHIJKLMNOPQRSTUVWXYZ!\"#$%&\'()*+,-./:;<=>?@[\\]^_`{|}~ \t\n\r\x0b\x0c'
```

Ist Reduktion von 128 auf 100 Zeichen wirklich den Aufwand wert?

Man verliert Umlaute!

Preprocessing

Vorverarbeiten von Text Daten.

Empfehlung: alles als Rohdaten ins Netz füttern.

Wort-basierte Tokenisierung

Wort-basierte Tokenisierung zerlegt den Text auf Wortebene.

„Das ist ein Beispiel.“ → ["Das", "ist", "ein", "Beispiel", "."]

Jedes Wort muss auf eine Zahl gemappt werden.

- **Vorteil:**

- Intuitiv und leicht verständlich.

- **Nachteile:**

- Empfindlich gegenüber **unbekannten** Wörtern (**Out-of-Vocabulary**, OOV z.B. compound-words!).
- Für neuronale Netze **unbrauchbar** weil Sprache keine abgeschlossene Menge von Worten ist
- (Zipf's Law)
- Problem: Worttrennung / Spaces!?

Used in **spacey**!? Haupt: Wort-basiert aber mit Erweiterungen für OOV Markierung für Trennzeichen

Tokenisierung

Tokenisierung

Mache aus Texten Listen von Vektoren. (kontextbasiertes Sprachdatenpaket)

Token = Zahl! => Vektor über **one-hot-encoding** als input

Character based [a, b, c] = [65, 66, 67] 100 / 128 ASCII 256 Byte Klassen / Zeichen

Wort-basierte Tokenisierung

„Das ist ein Beispiel.“ → ["Das", "ist", "ein", "Beispiel", "."] je nach Wörterbuch

Zahl = Zeile im Wörterbuch

Laut-basiert (Sprache) [Ha le lu ja] ≈ **Silben-basiert (Text)** Zahl mit mapping ['Ha'=>1 ...]

/ˈguːtən ˈmɔkŋn̩ ˈmaɪnə ˈʃɛːfçən/ geht über IPA hinaus.

One Hot Encoding

Generell: Mache aus Klassen Vektoren.

n Klassen => Vektor der Länge n

Indem man bei **Klasse i** an Stelle **i eine 1** setzt sonst alles null.

Subword Tokenisation

Subwort-Tokenisierung zerlegt den Text in *häufige* Subworteinheiten.

Beispiel: „unbelievable“, $\rightarrow$ ["un", "believ", "able", " ", ","].

"believe" $\Rightarrow$ ["believe"] oder ["believ", "e"] ? nicht: ["be", "l", "i", "e", "ve"]

Bei GPT:

bel
ieve

⚠ **semantisch oder nicht** (je nach Sprache aber auch innerhalb der Hauptsprache)

nicht notwendigerweise Inhaltsbasiert / Bedeutungstragend

Diese Methode wird oft bei Modellen wie BERT oder GPT verwendet, um ein **effizientes vollständiges** Vokabular zu schaffen.

Menge von Tokens = "Vokabular"

Vokabular enthält auch einzelne Zeichen "a", "b", ...

