

科学技術と倫理 「生成AIの原理」

日本語音声認識システムを 作るには？

「あ」を認識できるようにする

「い」を認識できるようにする

「う」を…

その通り！

このひらがな列を日本語かな漢字変換
のような技術で日本語文章へ



日本語の音素

1. 文
 2. 文
 3. 単
 4. 文
 5. ひ
-
4. 音

Eゲイト英和辞典

prefectural

音節 pre·fec·tur·al 発音記号・読み方

prɪfɛktʃ(ə)rəl

形容詞

(フランス・日本などの) 県の, 府の; 県立の, 府立の

- the prefectural office

県庁, 府庁

a, i, u, e, o, k, s, t, n, h, m, y, r, w, ɪ,
a:, i:, u:, e:, o:, g, z, d, gy, py,...



生成AIでは単語単位に扱われる

■ では、コンピュータ内での単語表現は？

■ まず、ビット(bit)から

- 2進数で表されます

- 1ビットは0か1の2通り

- 2ビットでは？

 - 最初の1ビットが01で次の1ビットが01

 - 00, 01, 10, 11 0～3の4通り 2×2

- 3ビットでは 000, ～ 111 の8通り $2 \times 2 \times 2$

1バイト(Byte)

1バイト = 8bit

- ・ $2^8 = 256$ 通り
- ・ 英文字や記号はこの256の番号で
(1バイトで)表している

Ex. “A” 65 0x41=0100 0001
64,32,16, 8,4,2

“a” 97 0x61=0110 0001

1バイト(Byte)=1英文字

“A” 65 0x41=0100 0001

“B” 66 0x42=0100 0010

“C” 67 0x43=0100 0011

“a” 97 0x61=0110 0001

“b” 98 0x62=0110 0010

“c” 99 0x63=0110 0011

漢字は？

2バイトで表される

- 1バイト: 256通りなので
- 2バイトにすると $256 \times 256 = 65,536$ 通り
- 65,536通りの漢字を表せる
- 日本語の漢字は、一般的に約2万字以上存在します (chat gptの答え)
- 常用漢字と呼ばれる漢字のリストに含まれる約2,136字 (chat gptの答え)

漢字コード

UTF8	UTF16	S-JIS	EUC-JP	JIS	文字	実体参照
E69F8E	67CE	9E74	DBD5	5B55	柎	26574;
E69F8F	67CF	9490	C7F0	4770	柏	26575;
E69F90	67D0	965E	CBBF	4B3F	某	26576;
E69F91	67D1	8AB9	B4BB	343B	柑	26577;

柏 10進数 26575

柏 2進数 0110,0111,1100,1111

0110 0111 **1100** 1111

8+4=12

柏 67**C**F

10→A 11→B 12→C 13→D 14→E 15→F

この漢字コードで何が分かりますか？

例えば

926a 1001 0010 0110 1010 男

8f97 1000 1111 1001 0111 女

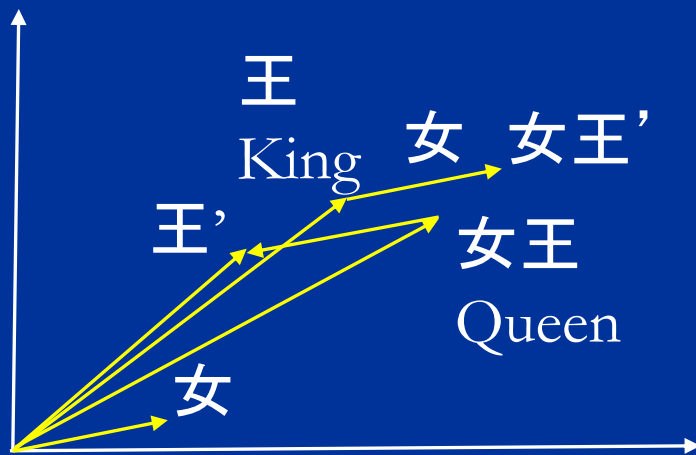
このような漢字コードでは
漢字の意味や単語の意味はとらえられない

単語と単語の近さや、関連度、共起関係などを
扱えるようにしたい

単語の分散表現 ベクトル化

単語の分散表現

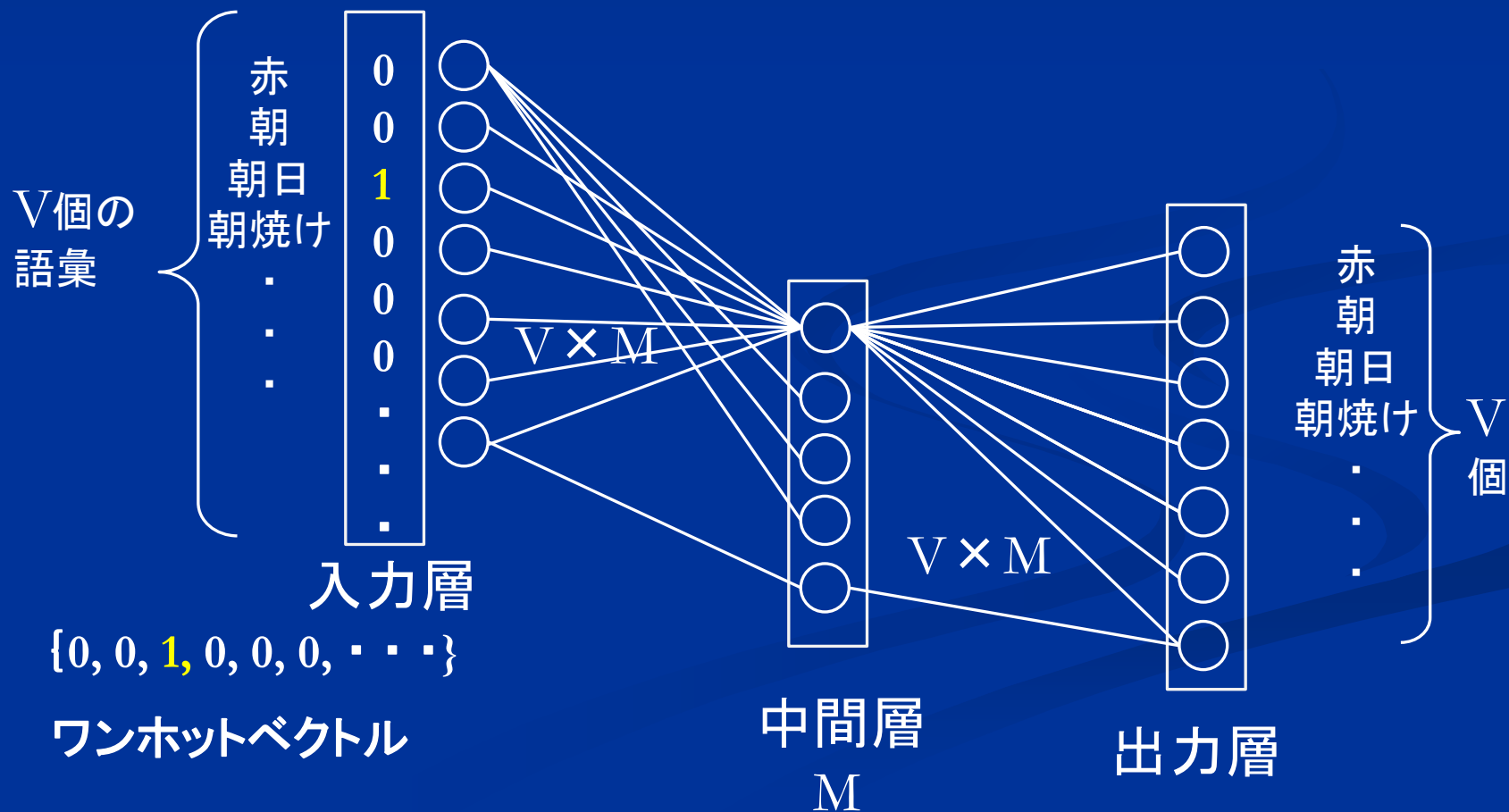
単語を分散表現(ベクトル)化
分散表現できると、単語の近さ、遠さが測れ、
合成ができそう



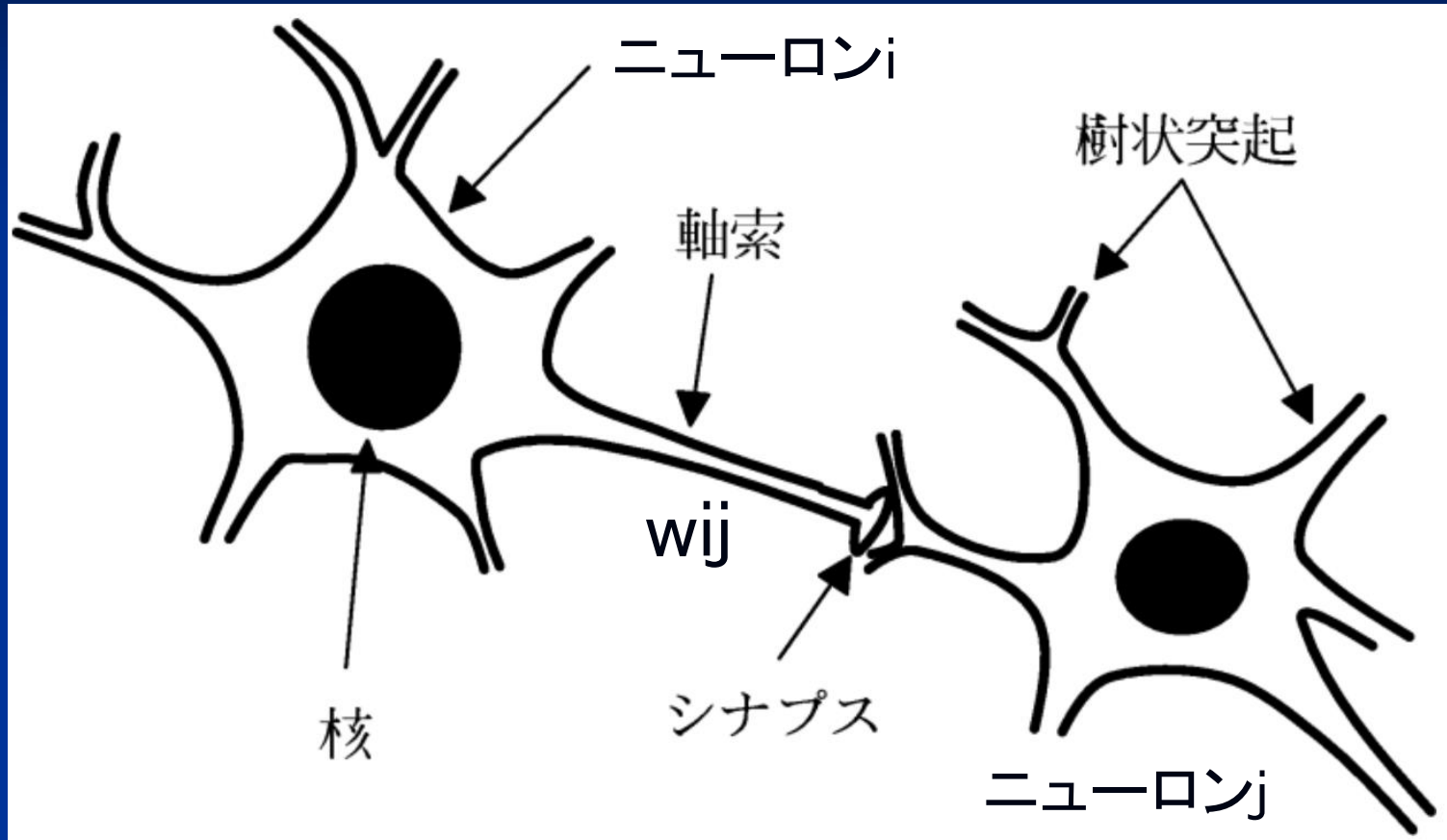
ニューラルネットによる分散表現

ニューラルネットへの入力 ワンホットベクトル

NNへ単語「朝日」を入力するには？



ニューロンと信号



- ニューロンiからニューロンjへの入力 x_{ij}
- そのシナプスの太さ(重み) w_{ij} $w_{ij} \times x_{ij}$

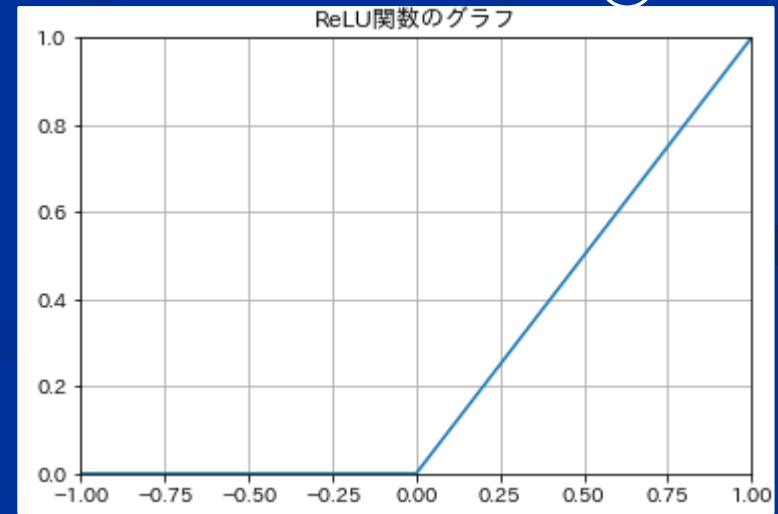
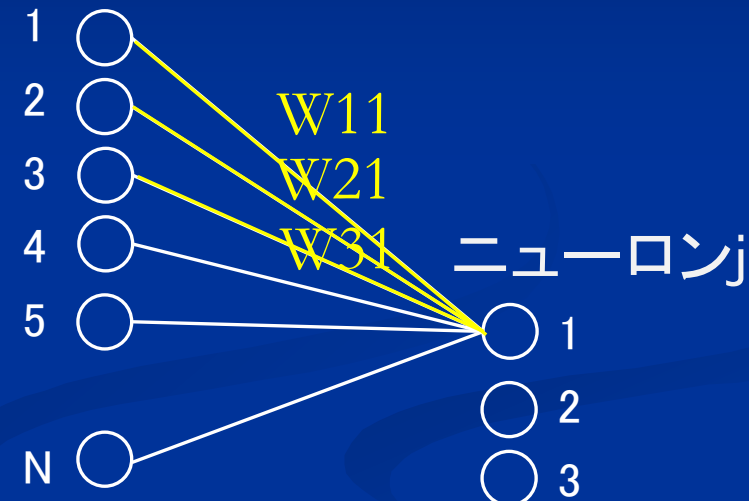
ニューロンjの入出力

- ニューロンjへはN個のニューロンから入力があるとすると

- ニューロンjへの入力は

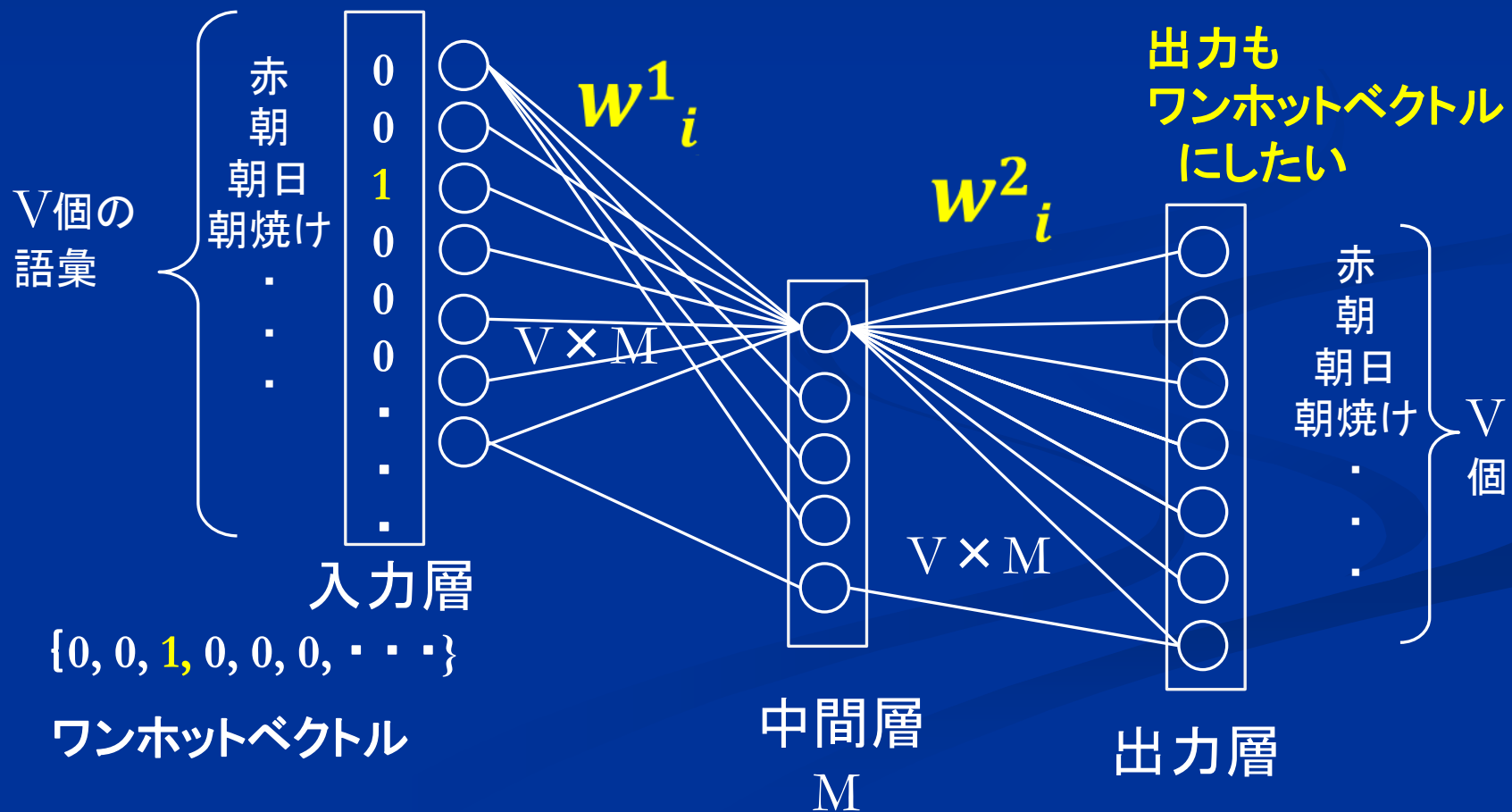
$$I_j = \sum_{i=0}^N W_{ij} \times x_i$$

- ニューロンjへの出力は、この入力に応じて



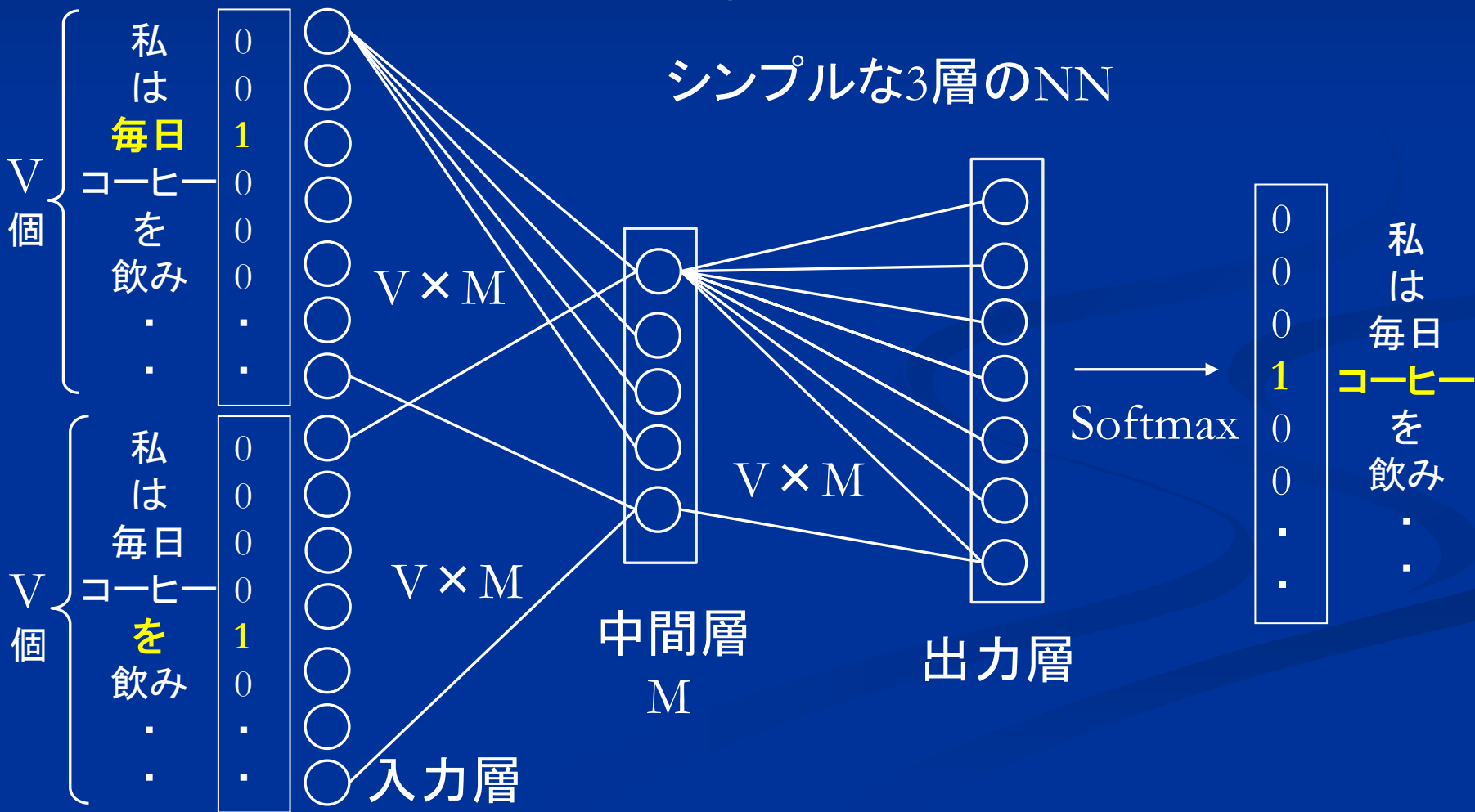
ニューラルネットへの入力 ワンホットベクトル

NNへ単語「朝日」を入力するには？



推論ベースの手法 word2vec

【？】にどのような単語が出現するのか推測



word2vec

1. word2vecモデル: CBOWとskip-gram

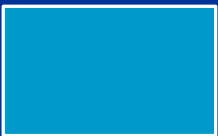
- CBOW (continuous bag-of-words) モデル: コンテキストからターゲットとなる単語を予想

私 は 毎日  を 飲み ます 。



$$L = -\frac{1}{T} \sum_{t=0}^T \log P(w_t | w_{t-2}, w_{t-1}, w_{t+1}, w_{t+2}) \quad L = -\frac{1}{T} \sum_{t=0}^T \log P(w_t | w_{t-1}, w_{t+1})$$

- skip-gramモデル: ターゲットから周囲の単語(コンテキスト)を推測

私 は  コーヒー  飲み ます 。



$$L = -\frac{1}{T} \sum_{t=0}^T \{\log P(w_{t-1} | w_t) + \log P(w_{t+1} | w_t)\}$$
$$L = -\frac{1}{T} \sum_{t=0}^T \{\log P(w_{t-1} | w_t) + \log P(w_{t+1} | w_t)\}$$

できるだけ正確な推測ができるように訓練⇒単語の分散表現

ニューラルネット言語モデルと 単語の分散表現

ニューラルネット言語モデル と 単語の分散表現

chat gpt の原理

大量のテキストデータから次の単語を予測する

$w_1 \ w_2 \ \dots \ w_{t-2} \ w_{t-1} \ w_t \ w_{t+1} \ w_{t+2} \ \dots \ w_T$

attention機構やtransformer機構など複雑なネットワークで構成され、単語の分散表現 $w'_1 \sim w'_{t-1}$ を入力していくと、次の単語 w'_t が予測できるように学習

$\underline{w_1 \ \dots \ w_{t-2} \ w_{t-1}} \quad w_t \ w_{t+1}$



Transformer

“Attention Is All You Need”で
2017年にGoogleから
発表されたモデル

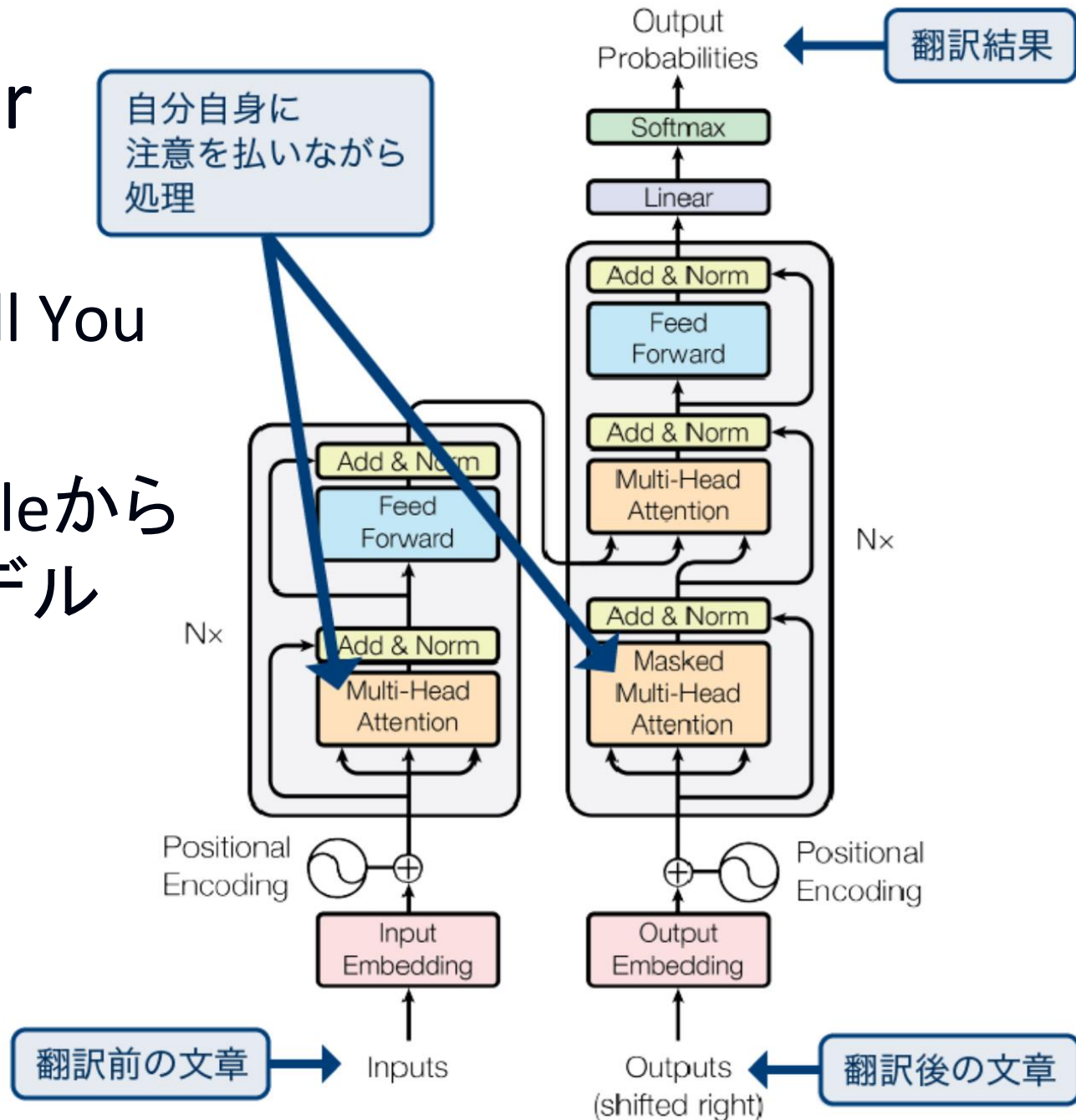


Figure 1: The Transformer - model architecture.

ChatGPTの出力

- 学習データ次第で決まる
- 望ましくないデータで学習すると。。。
- 望ましくないデータとは

ChatGPTの出力の利用

- ChatGPTの出力を使ってレポートを作成しました
- 不正行為になりますか？

レポートのコピー

- 友達のレポート見せてもらってそのままコピーを提出しました
 - 不正行為ですね？
 - どうなりますか？
-
- Reiki-Base インターネット版 体系目次
検索 (iwate-pu.ac.jp) の履修規定
 - (不正行為) 第10条

剽窃

- 他人の作品・学説などを自分のものとして発表すること
- 剽窃チェッカー

ChatGPTは使ってはいけないの？

- 本学では、まだ利用にあたっての明確なガイドラインがない（策定中）

東北大の例

- 確認してみてください。