

RoutesFormer: A sequence-based route choice Transformer for efficient path inference from sparse trajectories

Qiu, S., Qin, G., Wong, M., & Sun, J. (2024). RoutesFormer: A sequence-based route choice Transformer for efficient path inference from sparse trajectories. Transportation Research Part C: Emerging Technologies, 162, 104552.

2025/05/20 理論談話会#6

B4 水谷玲太



Abstract

1. 自然言語処理の深層学習モデルであるTransformerを応用し、**シーケンスベースの経路選択モデル RoutesFormer**を開発
2. 上海のタクシーデータセットで実験を行い、**既存のモデルを大きく上回る推定精度を記録**
3. モデル内のAttention機構の意義を**Ablation study**により考察し、経路選択行動のモデリングにおいて**過去の移動履歴を考慮する重要性**を提示

新規性・有用性・信頼性

- 新規性
 - 自然言語処理モデルであるTransformerの行動モデルへの応用
- 有用性
 - 疎な観測データの**経路選択**（Route Choice）と**経路補完**（Path Inference）を統一して扱うことが可能
 - マルコフ性仮定がなく、**過去の移動履歴の情報**を推論に用いることが可能
- 信頼性
 - 上海のタクシーデータを用いた実験で、既存の6つのモデルを全ての評価指標で上回る
 - アブレーション分析により、モデルの妥当性を検証した

目次

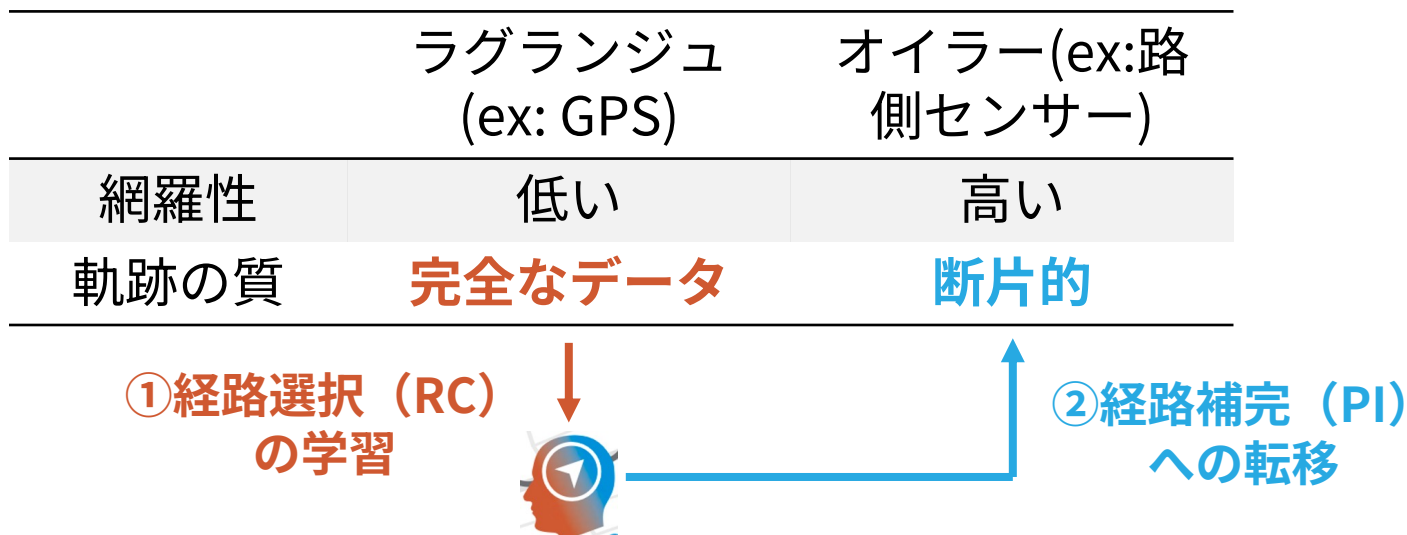
1. Introduction
2. Related work
 1. 既存の行動モデル
 2. Transformerの交通分野への応用例
3. Problem Description and Preliminaries
 1. 定義
 2. 自然言語処理の穴埋め問題とのアナロジー
4. Methodology
 1. Inputの下準備
 1. トークン化
 2. 埋め込み
 3. 位置エンコーディング
 2. エンコーダー・デコーダー
 1. Self-Attention
 2. Masked Self-Attention
 3. Cross-Attention
 4. Multi-Head Attention
 5. その他の構成
3. RoutesFormerにおけるAttentionの解釈
4. 補足事項
5. Case Study
 1. データ設定・比較対象
 2. 評価指標
 3. 推論精度の比較
 4. 計算時間の比較
 5. Ablation Study
 1. 精度の比較
 2. 精度の分布比較
 3. 推論結果の比較
6. Conclusion

1. Introduction

1.1 交通把握技術の進展とデータ統合の必要性

- センサー技術と機械学習の進展
 - ⇨交通システムの理解・管理・最適化はマクロ→ミクロレベルへ
- ⇨ミクロレベルの把握には**データの限界**:
 - **オイラー的観測** ex)GPS
 - 完全な軌跡を追えるが、サンプリング率が低い
 - **ラグランジュ的観測** ex)路側センサー
 - 全交通量を観測できるが、断片的な情報しか得られない
- ⇨ラグランジュ的データによる**①経路選択の学習** & オイラー的データの**②経路補完への転移**の枠組みが必要（次図）

1.1 交通把握技術の進展とデータ統合の必要性



⇒ 密な完全軌跡データを推論可能🙌

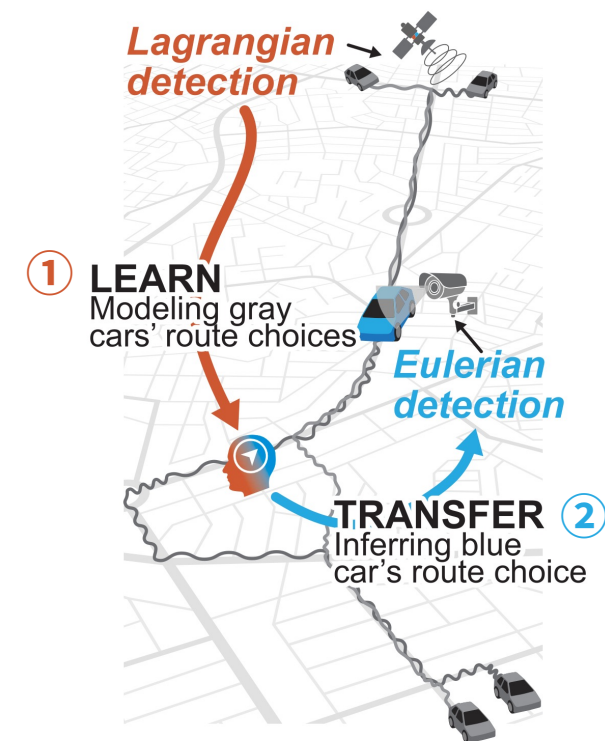
⇔ 既存のモデルの限界：

- マルコフ仮定…過去の移動履歴の影響無視
- 長い経路を複数の問題に分割 → 準最適に留まる



RoutesFormer：

- ①経路選択の学習と②経路補完の推論を統一的に処理
- 経路全体を一括で扱い、長期的な情報を学習



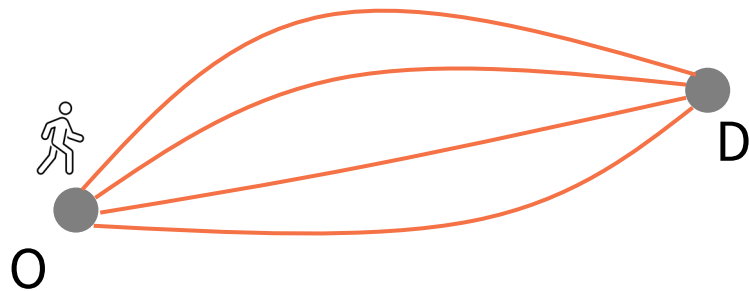
2. Related work

2.1 既存の行動モデル

①経路選択 (RC) モデル

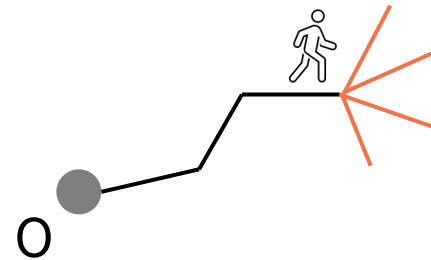
パスベース (パス列挙)

- PSL、DNN-PSL



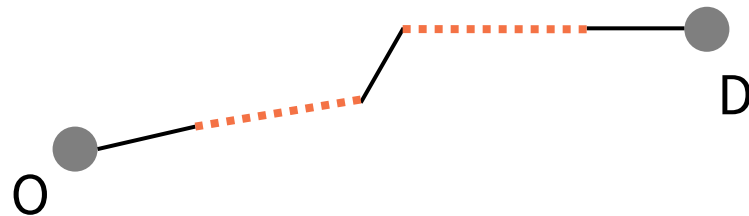
リンクベース (リンクを逐次選択)

- RL…マルコフ仮定



⇒ リンクからなる**系列全体を扱うシーケンススペース**のモデルへ (今回)

②経路補完 (PI) モデル



- 最短経路での補完→現実の選好との乖離
- DNNを用いた例も出始めたが、不連続パス全体の考慮ができない／複数のパス候補を独立して扱う⇨相関考慮なし・メモリ容量増大

2.2 Transformerの交通分野への応用例

- Transformer：従来のRNNと異なり、attention機構のみで完結する自然言語処理モデルとして成功を収める
- これまでの適用例：
 - Sequence prediction → 交通流予測、軌跡予測
 - Sequence to sequence → マップマッチング
- BERTのような空所補充課題への応用はない
- TransformerのSequence to sequence（翻訳）タスクと、BERTのcloze（穴埋め）タスクの組み合わせが、今回取り組む長さ可変の穴埋めタスク

⇒モデル設計（3～4章）

3. Problem Description and Preliminaries

3.1 定義

$G = (N, E)$: ネットワーク

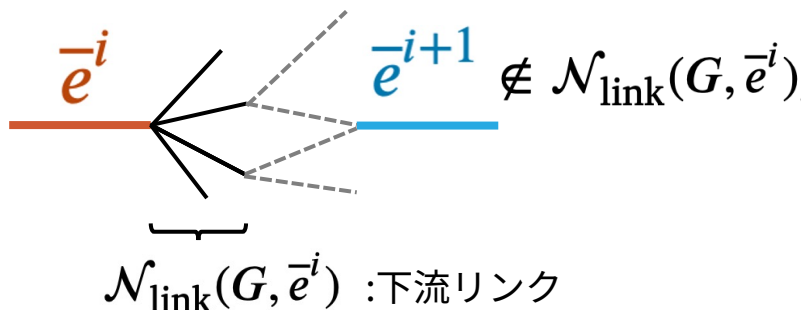
$\bar{R} = [\bar{e}^1, \bar{e}^2, \dots, \bar{e}^i, \dots, \bar{e}^{|\bar{R}|}]$: パス

$\mathcal{N}_{\text{link}}(G, \bar{e}^i)$: ネットワーク G におけるリンク $\bar{e}^i$ の下流リンクの集合

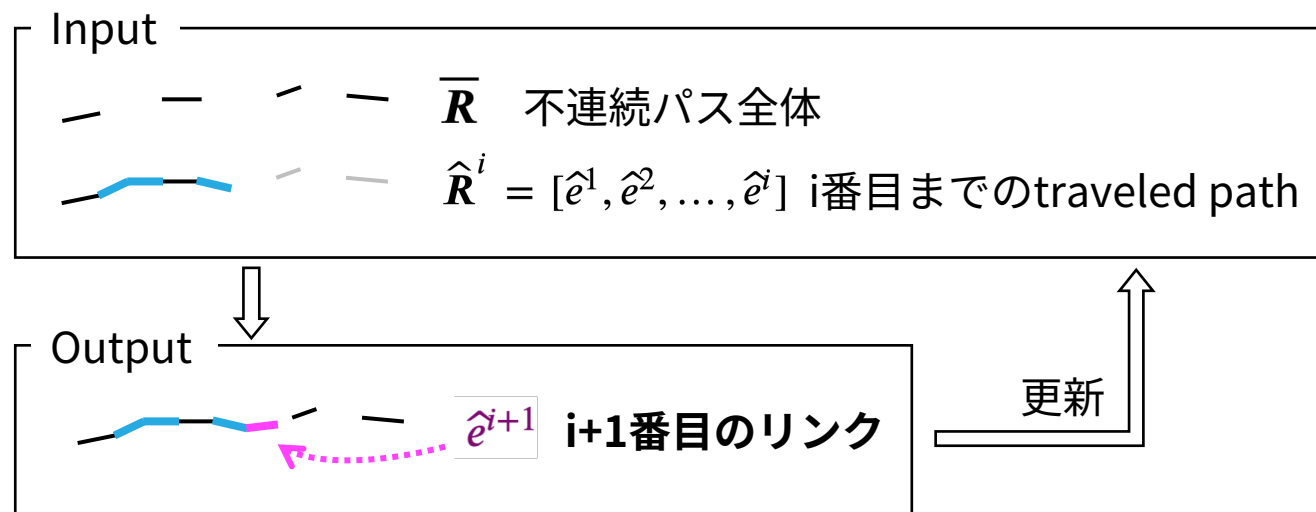


$\bar{e}^{i+1} \notin \mathcal{N}_{\text{link}}(G, \bar{e}^i)$ なる i がある時、

$\bar{R} = [\bar{e}^1, \bar{e}^2, \dots, \bar{e}^i, \dots, \bar{e}^{|\bar{R}|}]$ は不連続パス



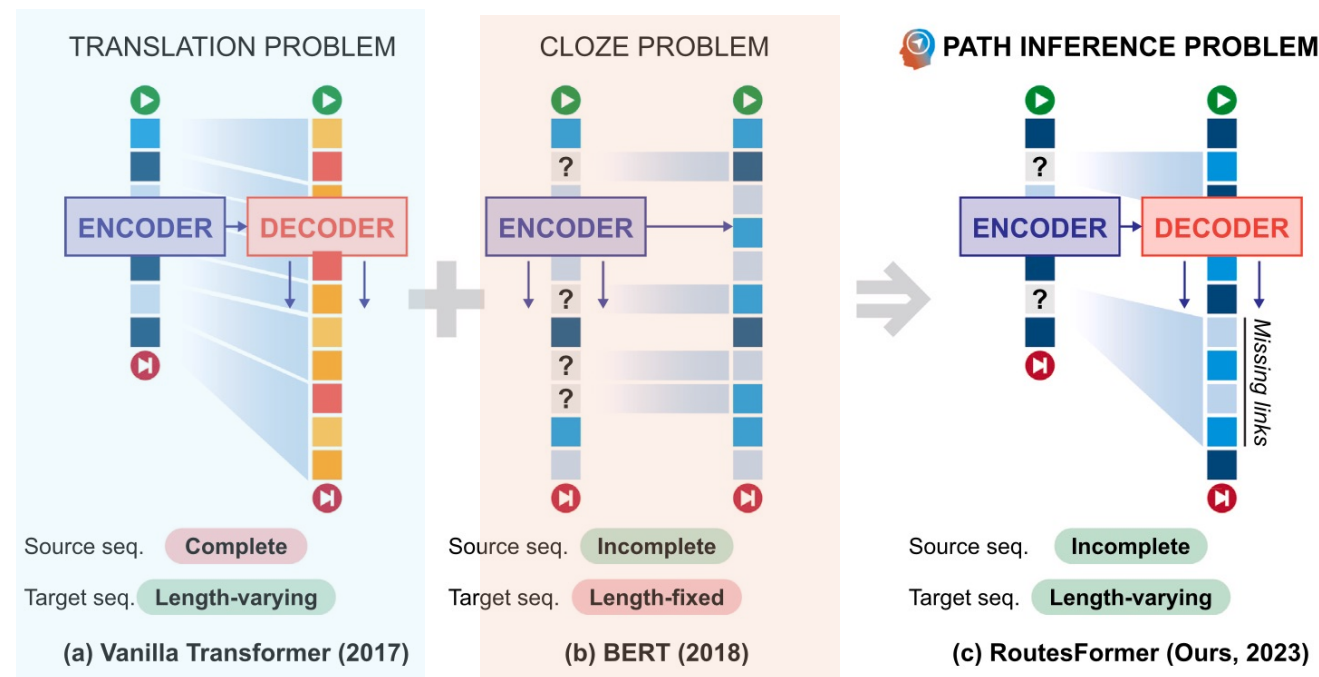
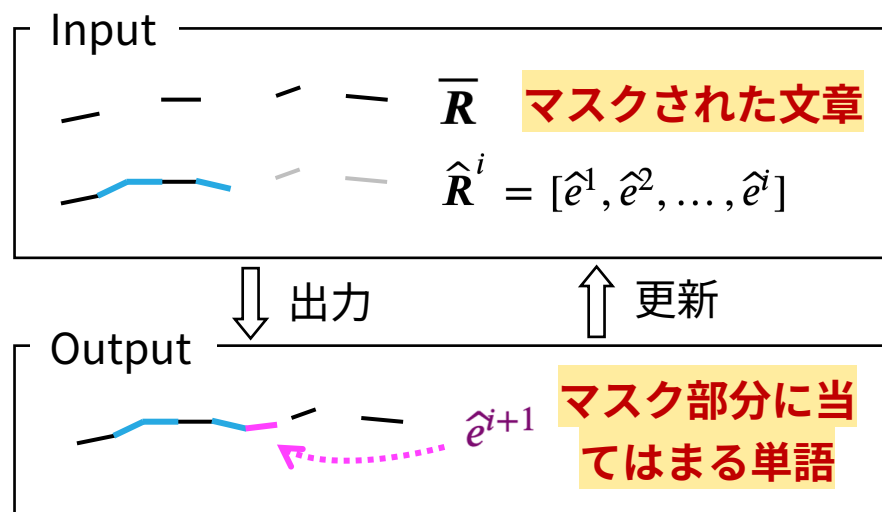
Sequence-based Route Choice(SRC) model



$$\min_{\hat{R}} \mathcal{L}(\bar{R}, \hat{R})$$

$$\text{s.t.} \quad \begin{cases} \hat{R} = [\hat{e}^1, \hat{e}^2, \hat{e}^3, \dots, \hat{e}^{|\hat{R}|}] \\ \bar{R} \text{ is the subsequence of } \hat{R} \\ \hat{e}^{i+1} = \text{SRC}(\bar{R}, \hat{R}^i), \quad \forall i \in \{1, \dots, |\hat{R}| - 1\} \\ \hat{e}^{i+1} \in \mathcal{N}_{\text{link}}(G, \hat{e}^i), \quad \forall i \in \{1, \dots, |\hat{R}| - 1\}. \end{cases}$$

3.2 自然言語処理の穴埋め問題とのアナロジー



本モデル特有の課題＝

空欄に当てはまるリンクの数が不明

⇔自然言語処理の空欄補充では空欄に入る単語数が既知

→出力の長さが可変である必要がある



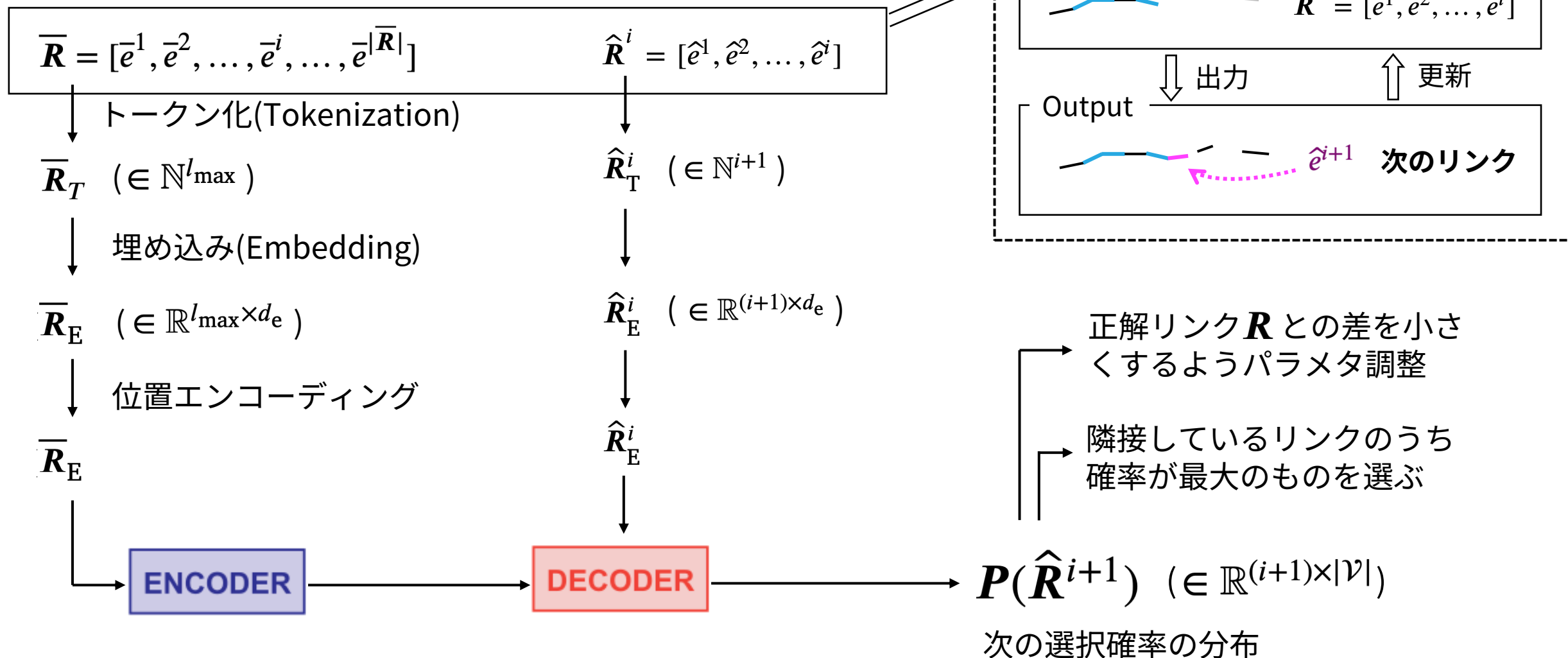
• Transformer (翻訳) の**Output長の可変性**

• BERT (空欄補充)

を組み合わせ、**不完全パスから完全パスに穴を埋めながら翻訳するモデル**

4. Methodology

4.1 inputの下準備



4.1.1 トークン化 (Tokenization)

- Inputをトークンのベクトルにする。種類は下記：
 - リンクIDトークン (自然数)
 - スペシャルトークン $\langle b \rangle \langle e \rangle \langle m \rangle \langle p \rangle$

Input

$$\begin{aligned} \overline{\mathbf{R}}_T &\in \mathbb{N}^{l_{\max}} : \langle b \rangle \ 1 \ 2 \ \langle m \rangle \ 5 \ 6 \ 7 \ 8 \ \langle m \rangle \ 10 \ \langle e \rangle \ \langle p \rangle \ \langle p \rangle \\ \hat{\mathbf{R}}_T^i &\in \mathbb{N}^{i+1} : \underbrace{\langle b \rangle \ 1 \ 2 \ 3 \ 4 \ 5}_{0 \rightarrow i} \end{aligned}$$

Traveled link

不連続パス。 $\langle e \rangle$ の後も
リンク最大長まで $\langle p \rangle$ で埋める

$$\begin{aligned} \mathbf{R}_T^{|R|} &\in \mathbb{N}^{|R|+1} : \langle b \rangle \ 1 \ 2 \ 3 \ 4 \ 5 \ 6 \ 7 \ 8 \ 9 \ 10 \\ \mathbf{R}_T^{|R|+1} &\in \mathbb{N}^{|R|+1} : \underbrace{1 \ 2 \ 3 \ 4 \ 5 \ 6 \ 7 \ 8 \ 9 \ 10 \ \langle e \rangle}_{1 \rightarrow |R|+1} \end{aligned}$$

左にずらす

7まで来た時の次の
推論の正解は8

推論時にtraveled linkとして
入れる正解リンク列

推論時に教師データとして入
れる正解リンク列

4.1.2 埋め込み (Embedding)

- トークン化したベクトルを、**リンク特徴量とともに**埋め込み次元ぶんのベクトルに
- 今回は埋め込み次元=32 (原文に128とあるが誤記?)

$$\text{Embedding}(e^i | \mathcal{W}_e) = \mathbf{W}_{ie}[e^i, :] \oplus e_f^{iT} \mathbf{W}_{fe}$$

学習可能な
パラメタ行列

リンク長や道路
の種類など
(Special tokenの
場合は0)

各トークンのベクトル

リンク e^i の特徴量ベクトル

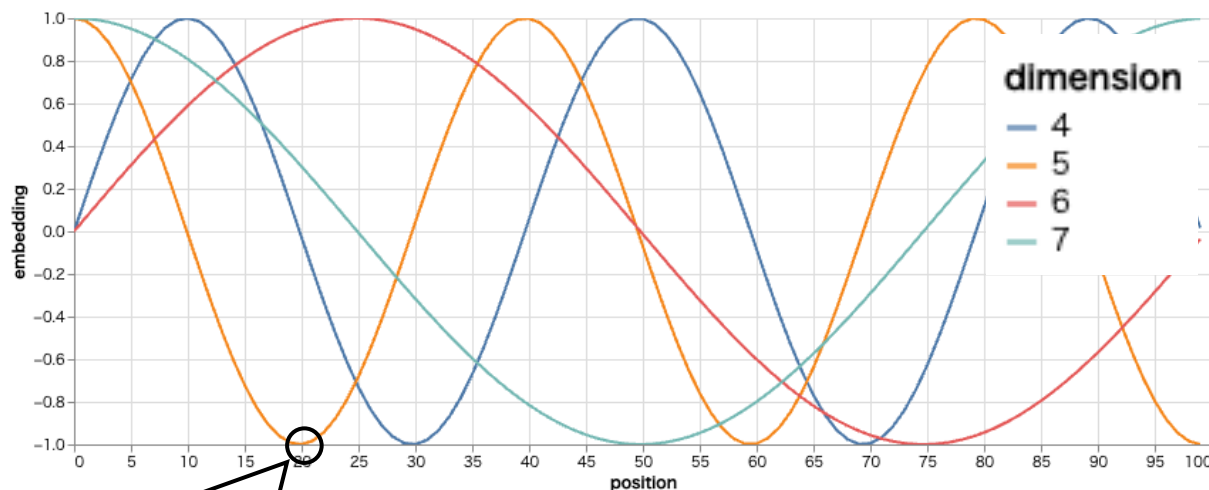
線形変換

$$= \left(\begin{array}{c} \text{トークン< b >のベクトル} \\ \text{トークン~~~のベクトル} \\ \dots \\ \text{リンク} e^i \text{のトークンのベクトル} \\ \dots \end{array} \right) \oplus \left(\begin{array}{c} \text{リンク} e^i \text{の特徴量ベクトル} \end{array} \right) \left(\begin{array}{c} \text{線形変換} \end{array} \right)$$
$$= \left(\begin{array}{cc} \text{リンク} e^i \text{のトークンのベクトル} & \text{線形変換されたリンク} e^i \text{の特徴量ベクトル} \end{array} \right)$$

4.1.3 位置エンコーディング (Positional Encoding)

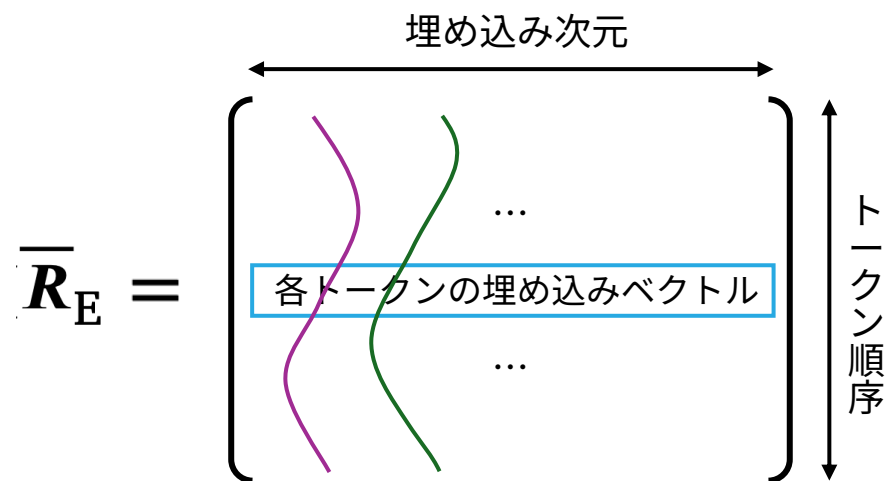
- この後のattention機構は系列の順序情報を考慮しない
- →順序をエンコーディングする

$$\begin{aligned} \text{PE}_{(t,2i)} &= \sin(t/10000^{2i/d_{\text{model}}}) \\ \text{PE}_{(t,2i+1)} &= \cos(t/10000^{2i/d_{\text{model}}}) \end{aligned} \quad \left(0 \leq i < \frac{d_{\text{model}}}{2}\right)$$



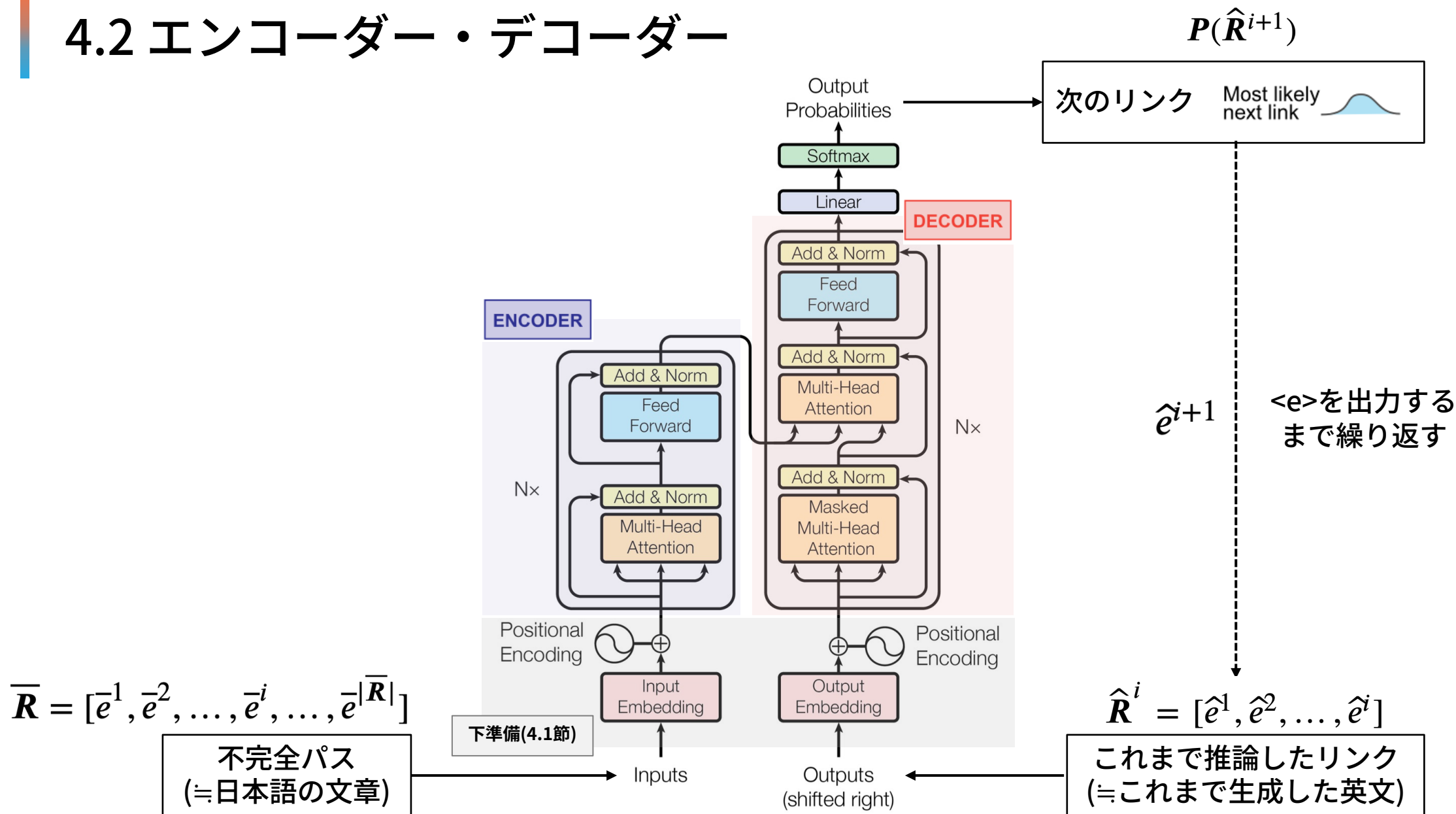
20番目のトークンの埋め込みベクトルの、5次元目に-1を足す

<https://nlp.seas.harvard.edu/annotated-transformer/>



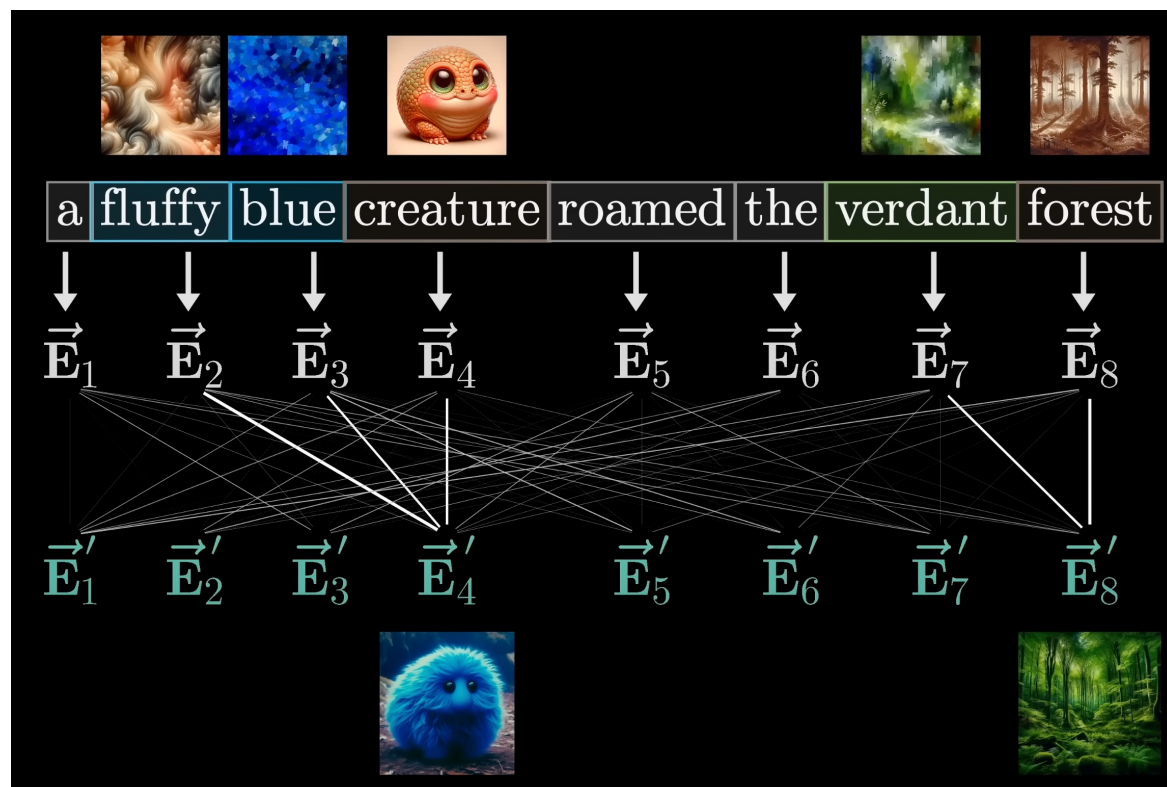
各次元ごとに色々な正弦波を足す

4.2 エンコーダー・デコーダー

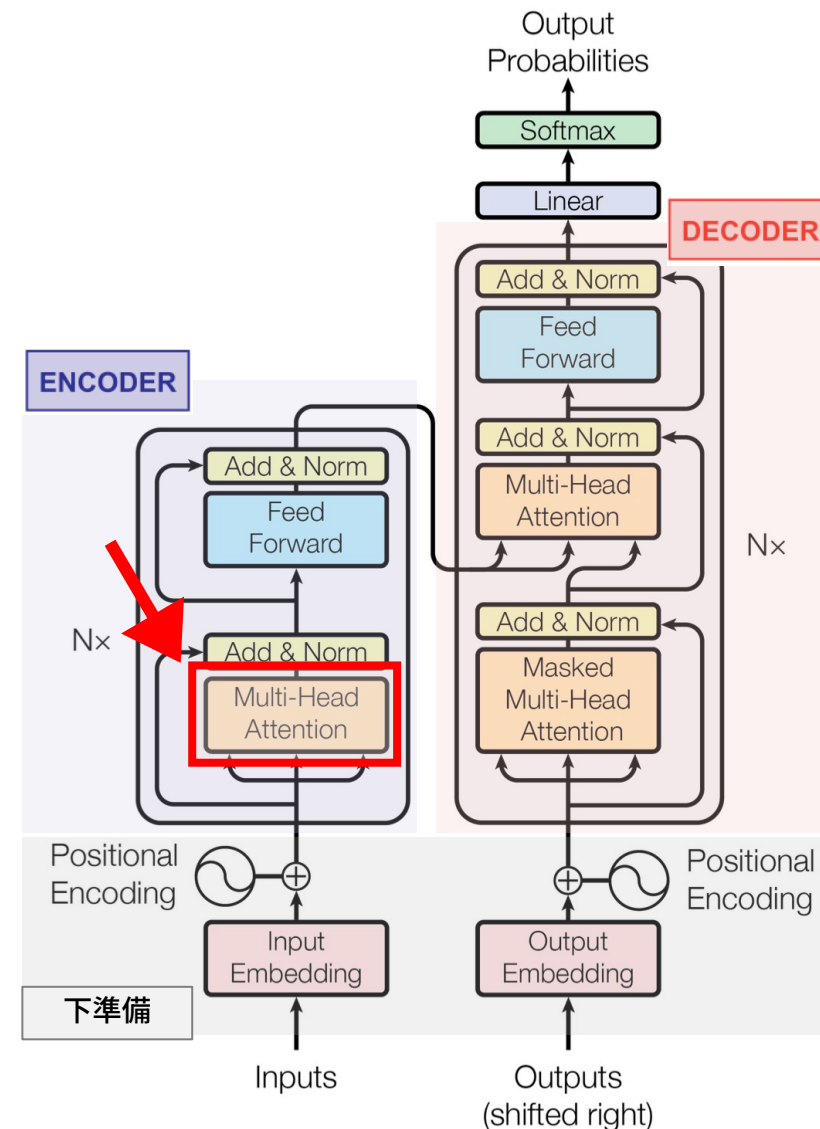
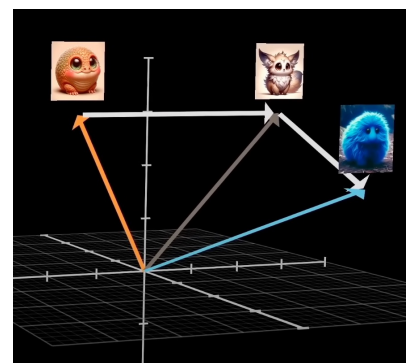


4.2.1 Self-Attention

- 文脈を加味していないinputの埋め込みベクトルを、文脈を考慮したベクトルに変換する



埋め込みベクトルの方向が意味を表すイメージ



以降添付している黒い背景の図：

https://www.youtube.com/watch?v=j3_VgCt18fA

4.2.1 Self-Attention

文脈を加味していないinputの埋め込みベクトルを、文脈を考慮したベクトルに変換する

- ①inputした埋め込みベクトル（今回は $X = Z$ で共通）の線形変換により Q, K, V を計算
- ② Q と K の各要素の積（を次元の平方根で割ったもの）のsoftmaxを計算し、 A とする
- ③ V の重みつき和（ $=A$ の要素との重み付き和）を出力する

$$1 \quad \mathbf{Q} \leftarrow \mathbf{XW}_q + \mathbf{1b}_q^T \quad [[\text{Query} \in \mathbb{R}^{l_x \times d_e}]]$$

$$2 \quad \mathbf{K} \leftarrow \mathbf{ZW}_k + \mathbf{1b}_k^T \quad [[\text{Key} \in \mathbb{R}^{l_z \times d_e}]]$$

$$3 \quad \mathbf{V} \leftarrow \mathbf{ZW}_v + \mathbf{1b}_v^T \quad [[\text{Value} \in \mathbb{R}^{l_z \times d_e}]]$$

$$4 \quad \mathbf{A} \leftarrow \text{softmax} \left(\mathbf{M} + \frac{\mathbf{QK}^T}{\sqrt{d_e}} \right) \quad [[\text{Attentionweight} \in \mathbb{R}^{l_x \times l_z}]]$$

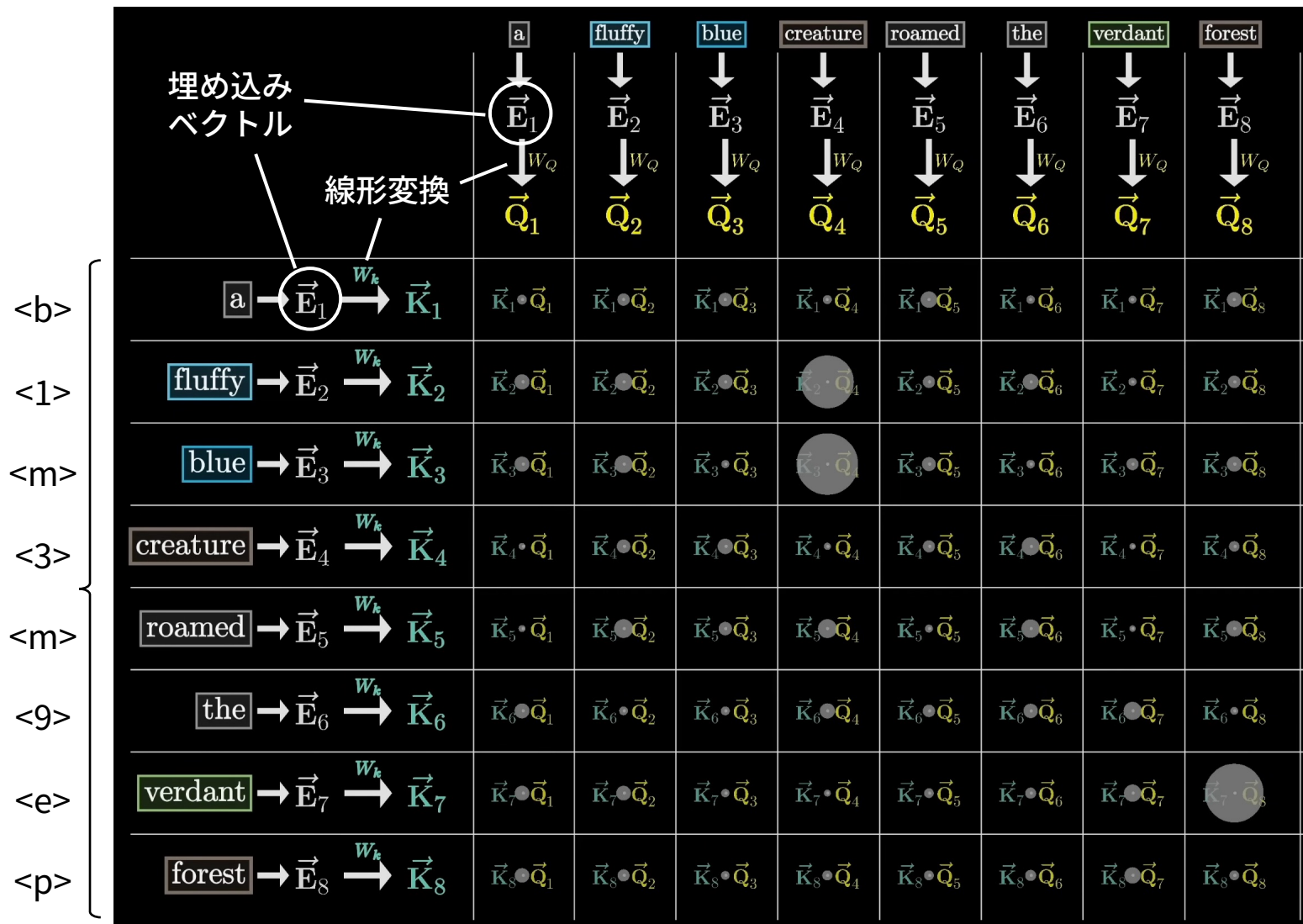
$$5 \quad \mathbf{Y} \leftarrow \mathbf{A} \cdot \mathbf{V}$$

4.2.1 Self-Attention

- ①inputした埋め込みベクトルの線形変換によりQ, K, Vを計算
- ②QとKの各要素の積（を次元の平方根で割ったもの）

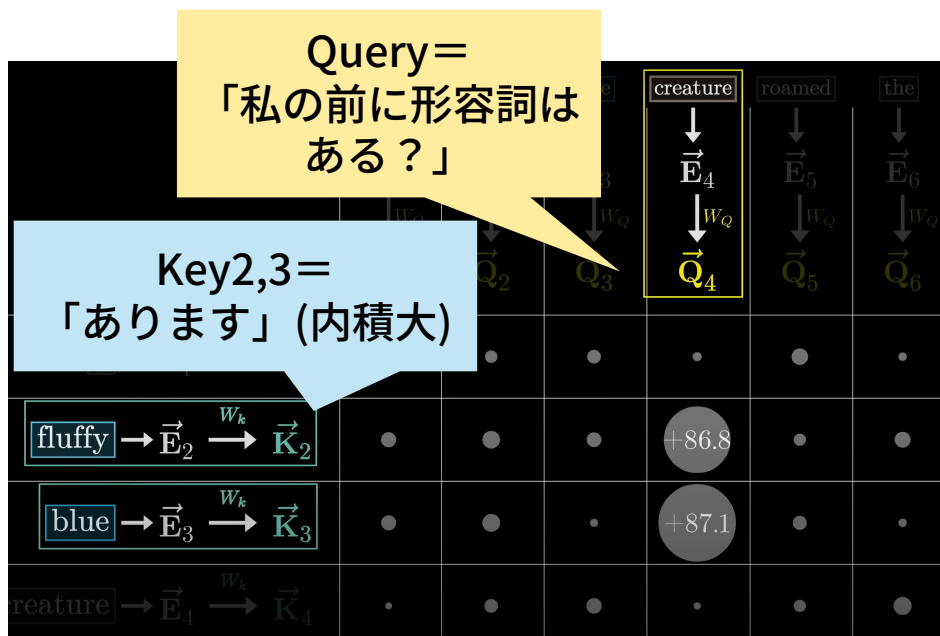
RoutesFormerでは下記のような不連続リンク列

 <1> <m> <3> <m> <9> <e> <p>



4.2.1 Self-Attention

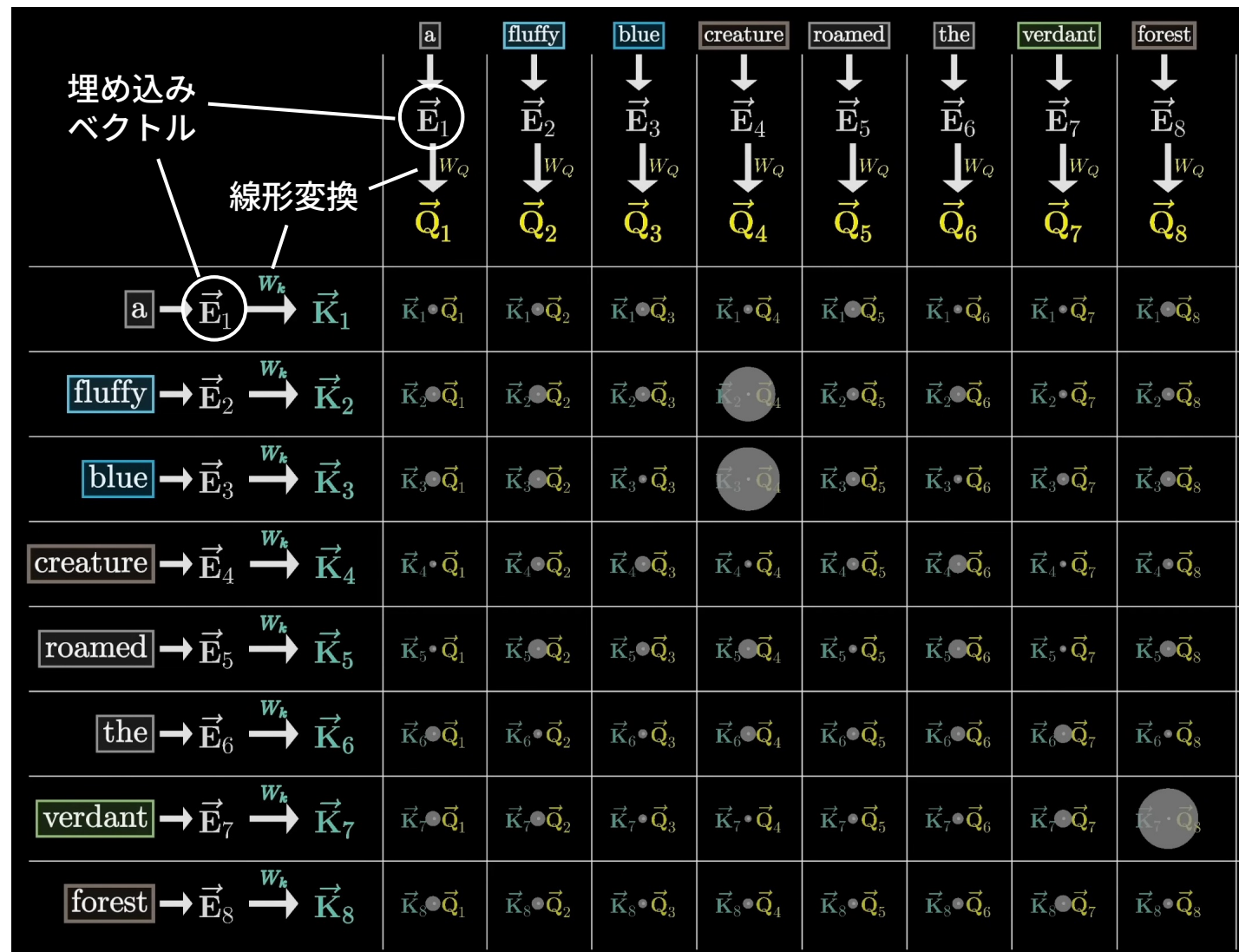
- ①inputした埋め込みベクトルの線形変換によりQ, K, Vを計算
- ②QとKの各要素の積（を次元の平方根で割ったもの）



↑ かなうように変換行列 W_Q, W_K の学習が進む

RoutesFormerでは下記のような不連続リンク列

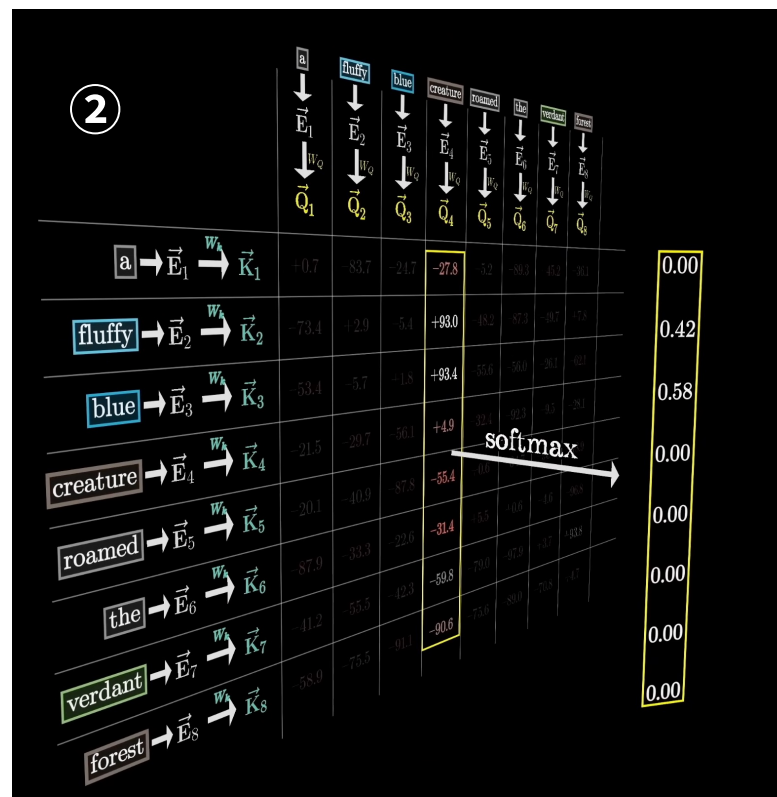
 <1> <m> <3> <m> <9> <e> <p>



4.2.1 Self-Attention

①inputした埋め込みベクトルの線形変換によりQ, K, Vを計算

②QとKの各要素の積（を次元の平方根で割ったもの）のsoftmaxを計算し、Aとする



Softmax関数：

$-\infty \sim +\infty$ の値の列を、 $0 \sim 1$ で和が1になるように変換する（→確率への変換）

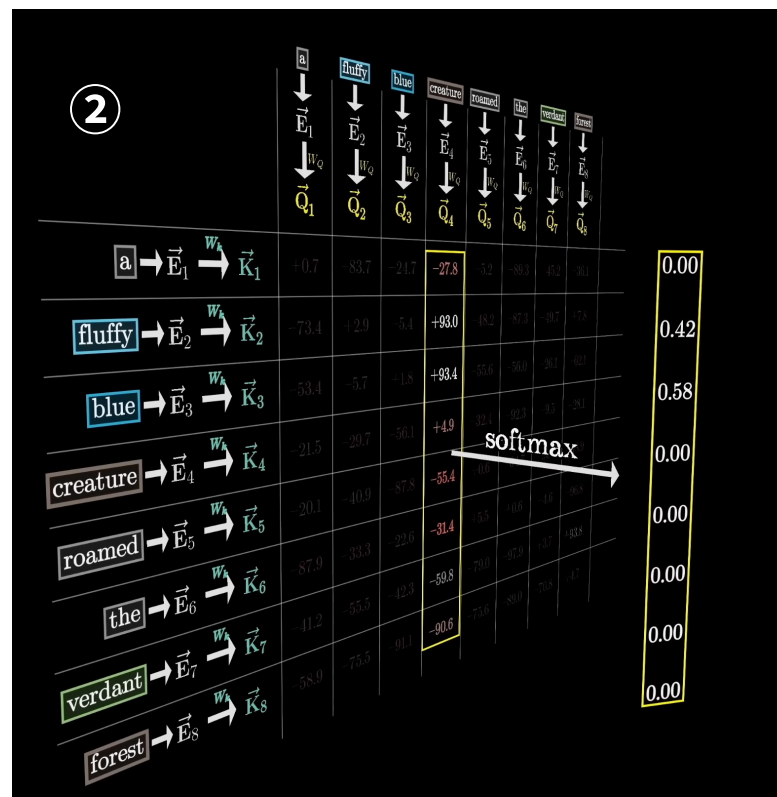
$$s(x_i) = \frac{e^{x_i}}{\sum_{j=1}^n e^{x_j}}$$

4.2.1 Self-Attention

①inputした埋め込みベクトルの線形変換によりQ, K, Vを計算

②QとKの各要素の積（を次元の平方根で割ったもの）のsoftmaxを計算し、Aとする

③Vの重みつき和（=Aの要素との重み付き和）を出力

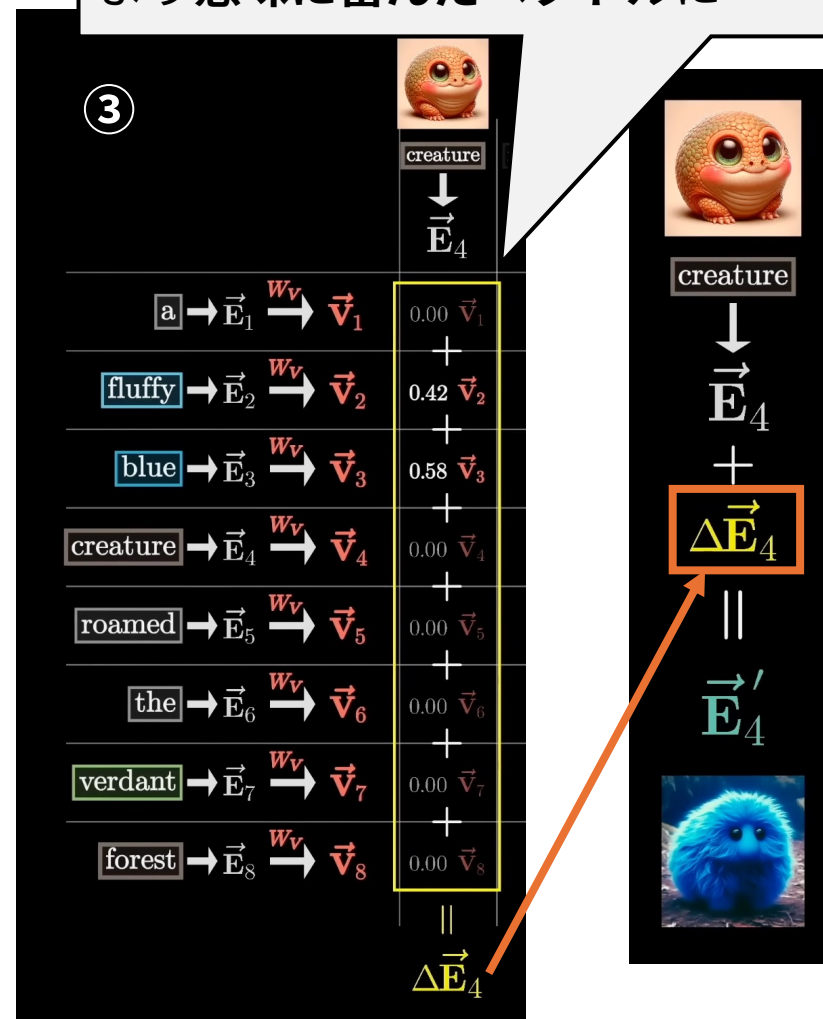


Softmax関数：

$-\infty \sim +\infty$ の値の列を、 $0 \sim 1$ で和が1になるように変換する（→確率への変換）

$$s(x_i) = \frac{e^{x_i}}{\sum_{j=1}^n e^{x_j}}$$

Creatureの埋め込みベクトルは、fluffyとblueのValueベクトルをAttentionの重みに応じて足し、より意味に富んだベクトルに



4.2.1 Self-Attention

①inputした埋め込みベクトルの線形変換によりQ, K, Vを計算

②QとKの各要素の積（を次元の平方根で割ったもの）のsoftmaxを計算し、Aとする

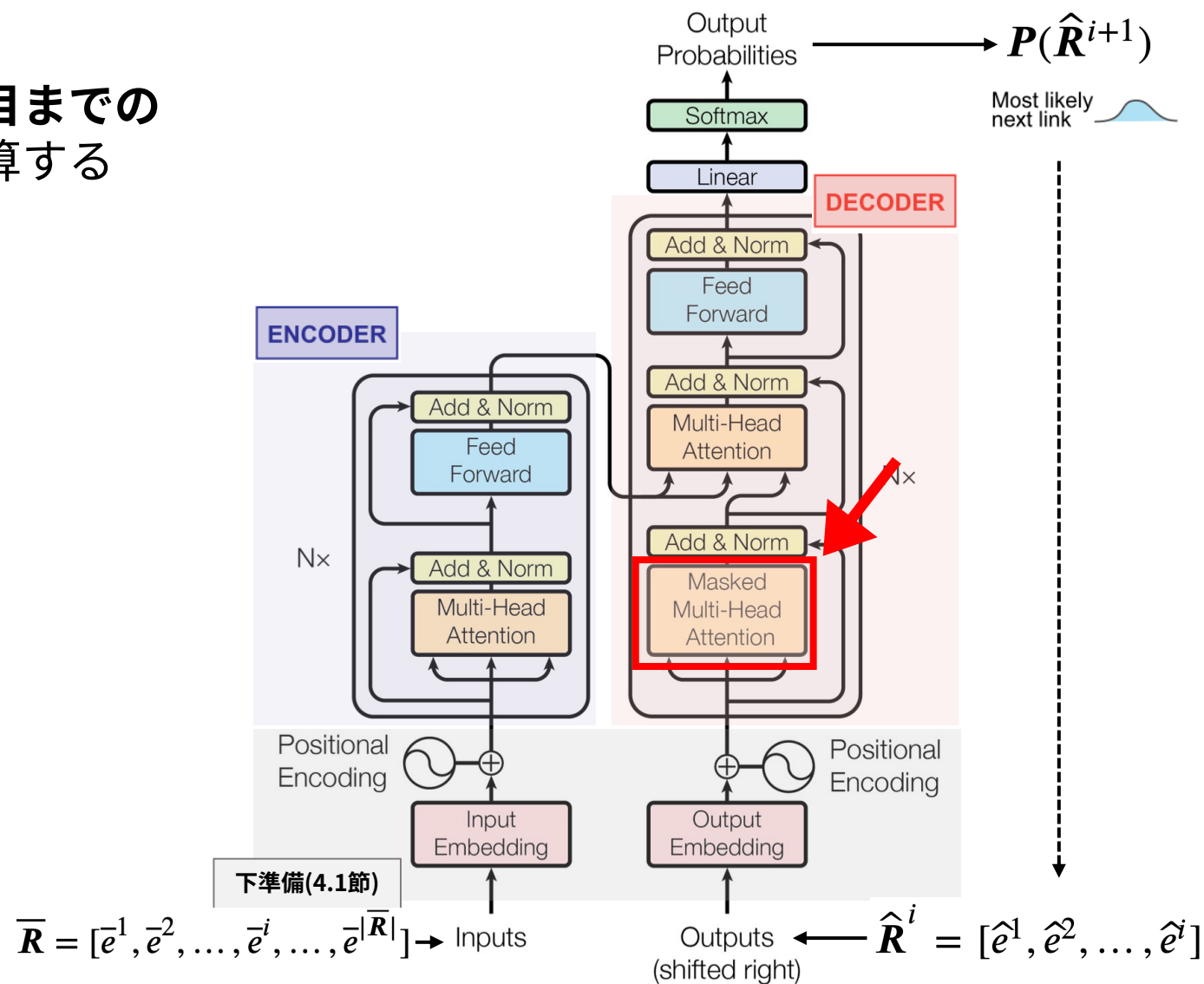
③Vの重みつき和（=Aの要素との重み付き和）を出力

⇒出力結果は、各トークンのベクトルを、より文脈を考慮した方向へ向けるベクトル

$$\begin{array}{ll} 1 & \mathbf{Q} \leftarrow \mathbf{XW}_q + \mathbf{1}b_q^T \quad [[\text{Query} \in \mathbb{R}^{l_x \times d_e}]] \\ 2 & \mathbf{K} \leftarrow \mathbf{ZW}_k + \mathbf{1}b_k^T \quad [[\text{Key} \in \mathbb{R}^{l_z \times d_e}]] \\ 3 & \mathbf{V} \leftarrow \mathbf{ZW}_v + \mathbf{1}b_v^T \quad [[\text{Value} \in \mathbb{R}^{l_z \times d_e}]] \\ 4 & \mathbf{A} \leftarrow \text{softmax} \left(\mathbf{M} + \frac{\mathbf{QK}^T}{\sqrt{d_e}} \right) \quad [[\text{Attentionweight} \in \mathbb{R}^{l_x \times l_z}]] \\ 5 & \mathbf{Y} \leftarrow \mathbf{A} \cdot \mathbf{V} \end{array}$$

4.2.2 Masked Self-Attention (計算上の処理)

- $i+1$ 番目の推論時、デコーダーでは i 番目までの traveled linksのSelf-Attentionを計算する



4.2.2 Masked Self-Attention (計算上の処理)

- $i+1$ 番目の推論時、デコーダーでは i 番目までの traveled linksのSelf-Attentionを計算する

⇨学習時には、完全なパス $R_T^{[R]}$ をinputするが、
学習時に $i+1$ 番目以降の正解を見て学習してほしくない

- ⇨ $i+1$ 番目以降のリンクへのattentionは $-\infty$ とする (マスクする)

→softmaxで0になる

Complete path $R_T^{[R]}$

<1>

<2>

<3>

<4>

<5>

<1>

<2>

<3>

<4>

<5>

+3.53	+0.80	+1.96	+4.48	+3.74	-
$-\infty$	-0.30	-0.21	+0.82	+0.29	-
$-\infty$	$-\infty$	+0.89	+0.67	+2.99	-0.41
$-\infty$	$-\infty$	$-\infty$	+1.31	+1.73	-1.48
$-\infty$	$-\infty$	$-\infty$	$-\infty$	+3.07	+2.94
$-\infty$	$-\infty$	$-\infty$	$-\infty$	$-\infty$	+0.31

Normalized Attention Pattern

<1>

<2>

<3>

<4>

<5>

<1>

<2>

<3>

<4>

<5>

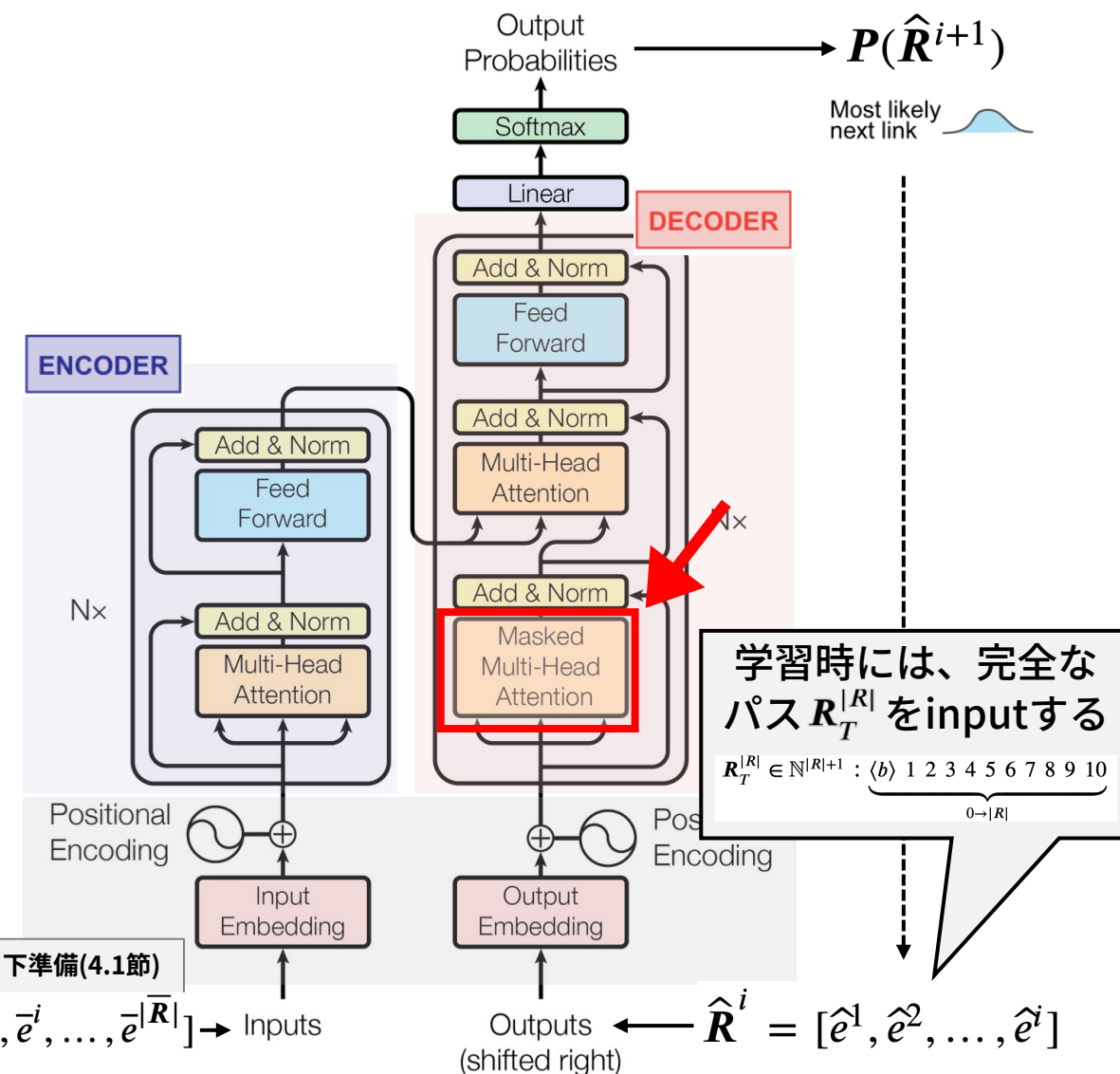
0.69	0.92	0.46	0.00		
0.08	0.02	0.01	0.46		
0.00	0.00	0.24	0.02	0.22	0.02
0.00	0.00	0.00	0.04	0.06	0.01
0.00	0.00	0.00	0.00	0.24	0.48
0.00	0.00	0.00	0.00	0.00	0.03

<5>の推論時には、~<4>までを参照する

softmax

<5>の推論時には、
~<4>までを参照する

softmax



4.2.2 Masked Self-Attention (計算上の処理)

- $i+1$ 番目の推論時、デコーダーでは i 番目までの traveled linksのSelf-Attentionを計算する

⇨学習時には、完全なパス $R_T^{[R]}$ をinputするが、
学習時に $i+1$ 番目以降の正解を見て学習してほしくない

- ⇨ $i+1$ 番目以降のリンクへのattentionは $-\infty$ とする (マスクする)
→softmaxで0になる

Complete path $R_T^{[R]}$

<1>

<2>

<3>

<4>

<5>

<1>

<2>

<3>

<4>

<5>

+3.53

+0.80

+1.96

+4.48

+3.74

-0.31

$-\infty$

-0.30

-0.21

+0.82

+0.29

+0.31

$-\infty$

$-\infty$

+0.89

+0.67

+2.99

-0.41

$-\infty$

$-\infty$

$-\infty$

+1.31

+1.73

-1.48

$-\infty$

$-\infty$

$-\infty$

$-\infty$

+3.07

+2.94

$-\infty$

$-\infty$

$-\infty$

$-\infty$

$-\infty$

+0.31

<5>の推論時には、
~<4>までを参照する

softmax

0.69

0.92

0.46

0.00

0.08

0.02

0.01

0.46

0.00

0.00

0.24

0.02

0.22

0.02

0.00

0.00

0.00

0.04

0.06

0.01

0.00

0.00

0.00

0.00

0.24

0.48

0.00

0.00

0.00

0.00

0.00

0.03

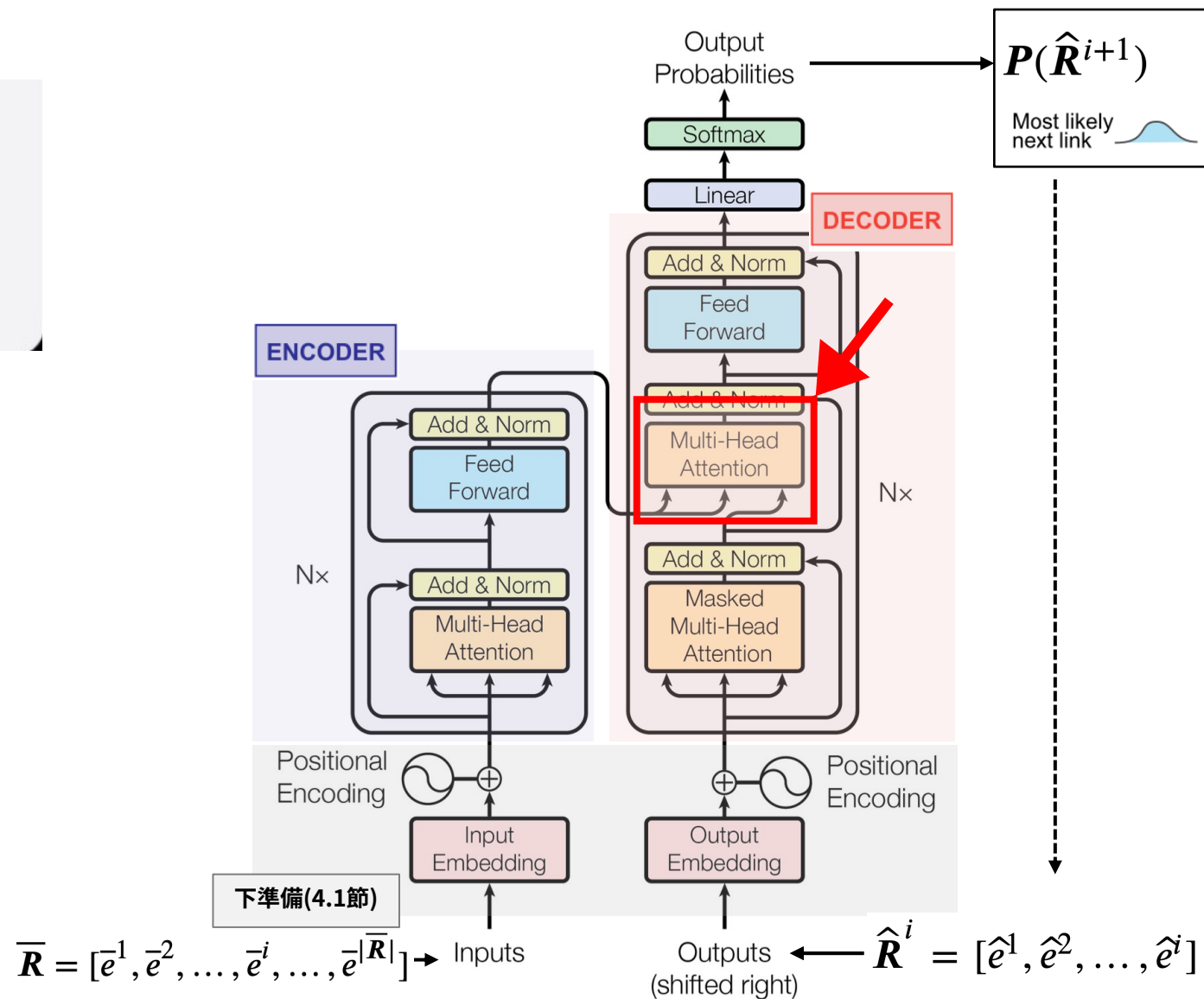
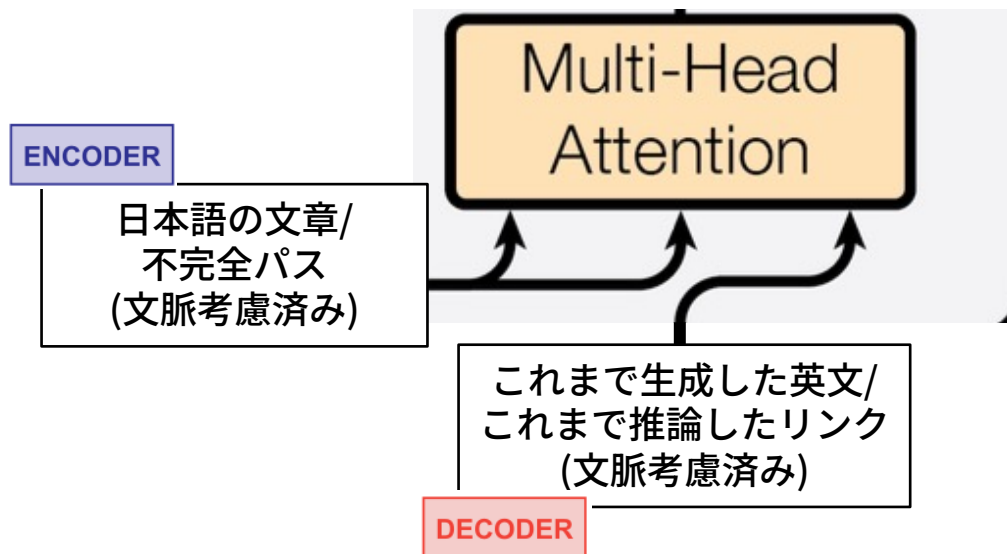
<5>の推論時には、
~<4>までを参照する

softmax

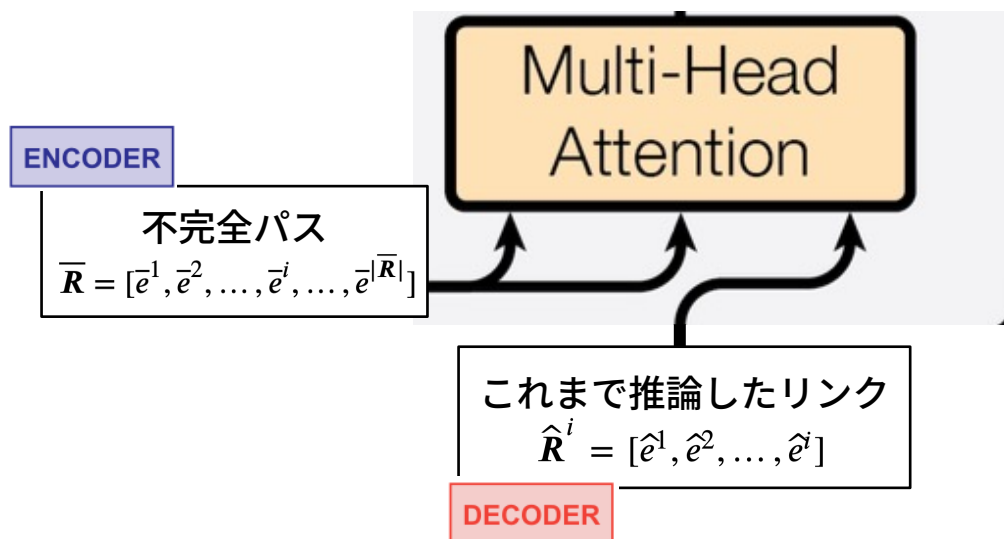
$$\begin{aligned}
 1 \quad & \mathbf{Q} \leftarrow \mathbf{X}\mathbf{W}_q + \mathbf{1}\mathbf{b}_q^T \quad [[\text{Query} \in \mathbb{R}^{l_x \times d_e}]] \\
 2 \quad & \mathbf{K} \leftarrow \mathbf{Z}\mathbf{W}_k + \mathbf{1}\mathbf{b}_k^T \quad [[\text{Key} \in \mathbb{R}^{l_z \times d_e}]] \\
 3 \quad & \mathbf{V} \leftarrow \mathbf{Z}\mathbf{W}_v + \mathbf{1}\mathbf{b}_v^T \quad [[\text{Value} \in \mathbb{R}^{l_z \times d_e}]] \\
 4 \quad & \mathbf{A} \leftarrow \text{softmax} \left(\mathbf{M} + \frac{\mathbf{Q}\mathbf{K}^T}{\sqrt{d_e}} \right) \quad [[\text{Attentionweight} \in \mathbb{R}^{l_x \times l_z}]] \\
 5 \quad & \mathbf{Y} \leftarrow \mathbf{A} \cdot \mathbf{V}
 \end{aligned}$$

$$\mathbf{M}_{\mathbf{U}}[i_x, i_z] = \begin{cases} 0, & i_z \leq i_x, \\ -\infty, & \text{otherwise.} \end{cases}$$

4.2.3 Cross-Attention

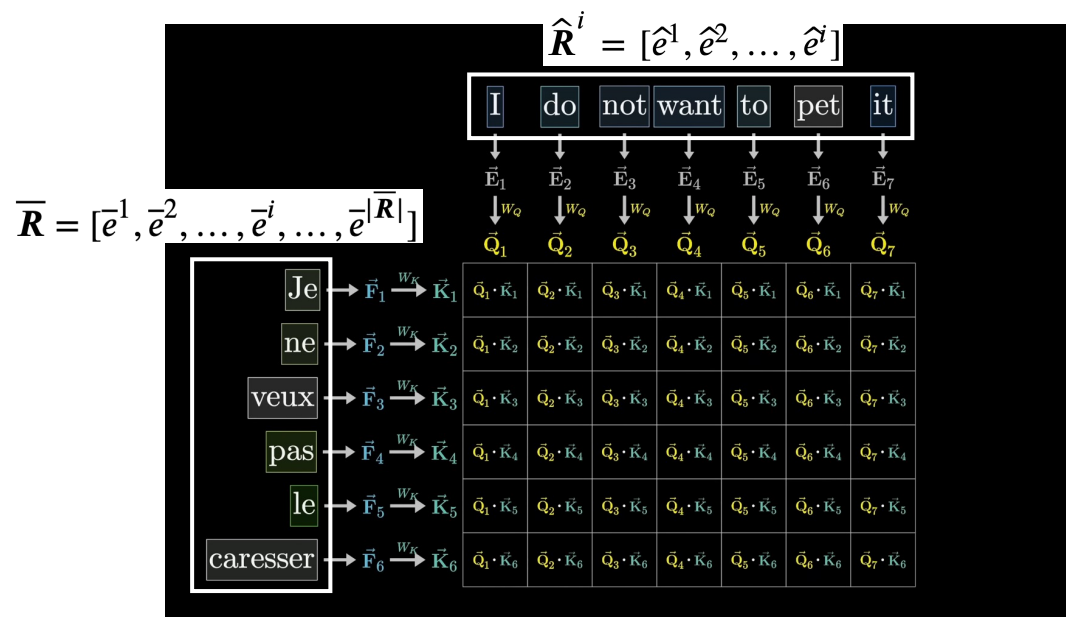


4.2.3 Cross-Attention



- 1 $Q \leftarrow XW_q + \mathbf{1}b_q^T$ $[[\text{Query} \in \mathbb{R}^{l_x \times d_e}]]$
- 2 $K \leftarrow ZW_k + \mathbf{1}b_k^T$ $[[\text{Key} \in \mathbb{R}^{l_z \times d_e}]]$
- 3 $V \leftarrow ZW_v + \mathbf{1}b_v^T$ $[[\text{Value} \in \mathbb{R}^{l_z \times d_e}]]$
- 4 $A \leftarrow \text{softmax} \left(M + \frac{QK^T}{\sqrt{d_e}} \right)$ $[[\text{Attentionweight} \in \mathbb{R}^{l_x \times l_z}]]$
- 5 $Y \leftarrow A \cdot V$

- 推論時に、 $\bar{R}$ のうちのリンクに注目すべきか



4.2.4 Multi-Head Attention

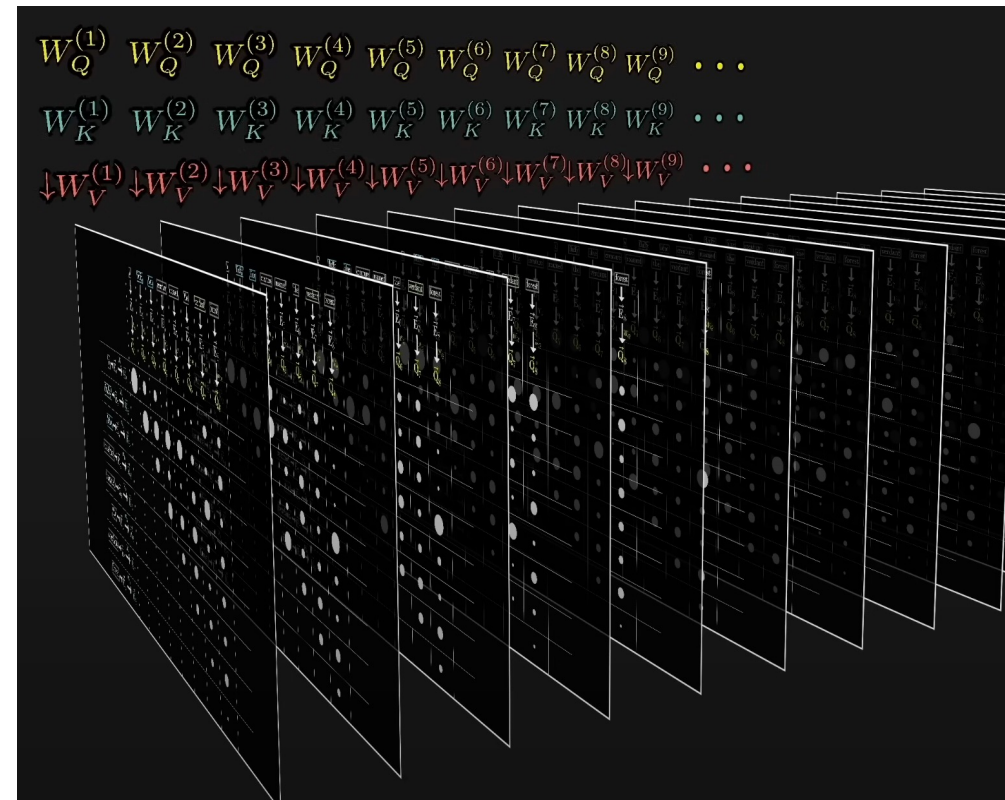
- 実際は何種類ものQ,K,Vの組み合わせを作り、**それぞれの組み合わせについてのAttentionとVの重み付き和を計算し、それらを結合した行列を出力する**

$$Y^h = \text{Attention}(X, Z | \mathcal{W}_{qkv}^h, \mathbf{M}),$$

$$\text{MHAttention}(X, Z | \mathcal{W}, \mathbf{M}) = [Y^1, \dots, Y^h, \dots, Y^H] \mathbf{W}_o + \mathbf{1} \mathbf{b}_o^T,$$

⇒要素間の多様な表現（文法構造、意味構造…）を抽出

RoutesFormerではヘッド数8



4.2.5 その他の構成



線形変換→ReLU→線形変換

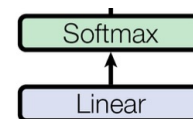
$$\text{FFN}(X|W_1, b_1, W_2, b_2) = \text{ReLU}(XW_1 + \mathbf{1}b_1^T)W_2 + \mathbf{1}b_2^T$$

(32次元 (埋め込み次元) → 中間層128次元 → 32次元)



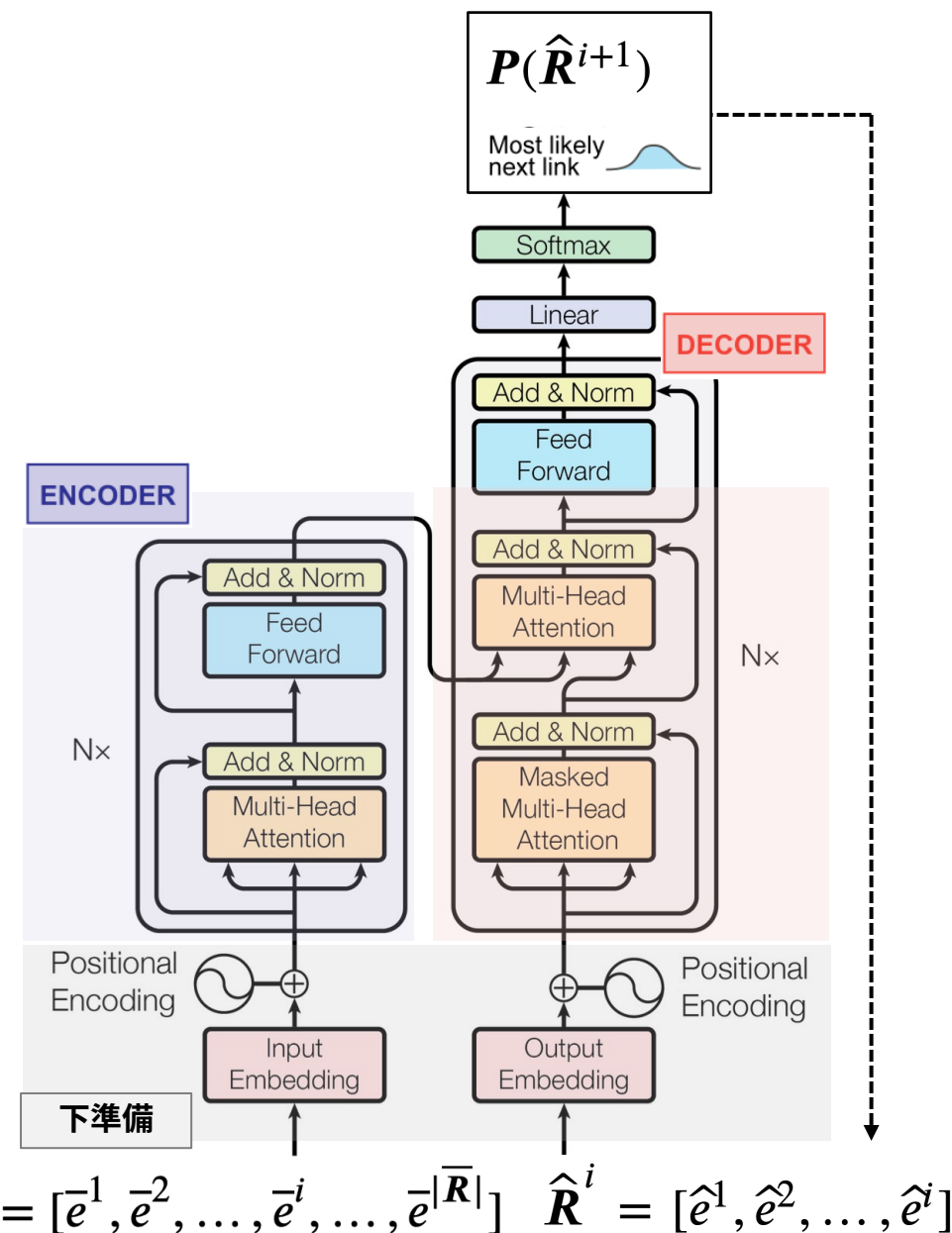
ネットワークを通す前の行列にネットワークの出力を足し
(残差結合)、安定化のために正規化

$$\text{LN}(X^i|\gamma, \beta) = \gamma \frac{X^i - \mu}{\sqrt{\sigma^2 + \epsilon}} + \beta$$

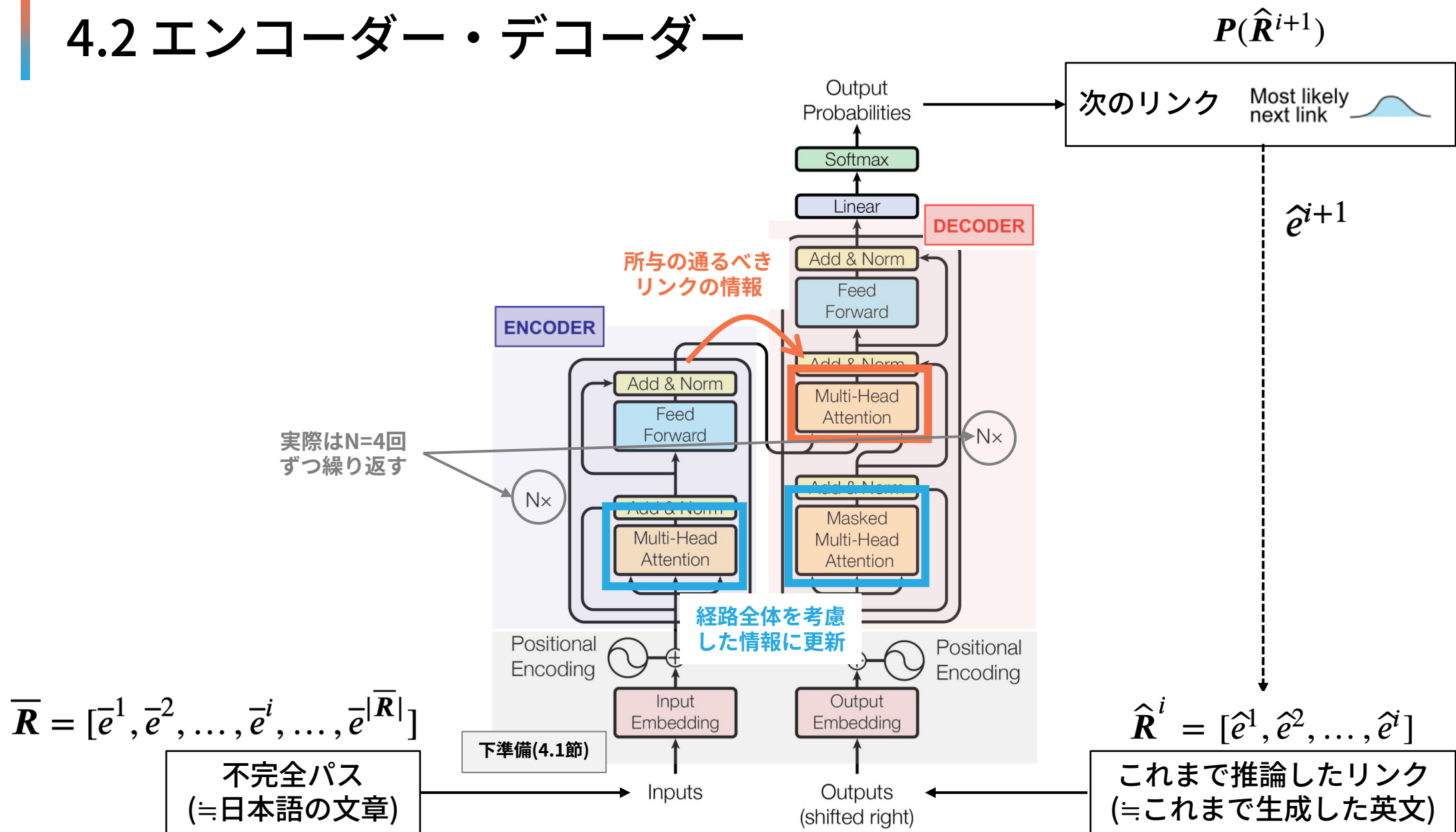


出力を線形変換したのちに、softmax
→次の単語/リンクの選択確率として出力

$$\text{Unembedding}(X|W_{\text{ue}}) = \text{softmax}(XW_{\text{ue}})$$



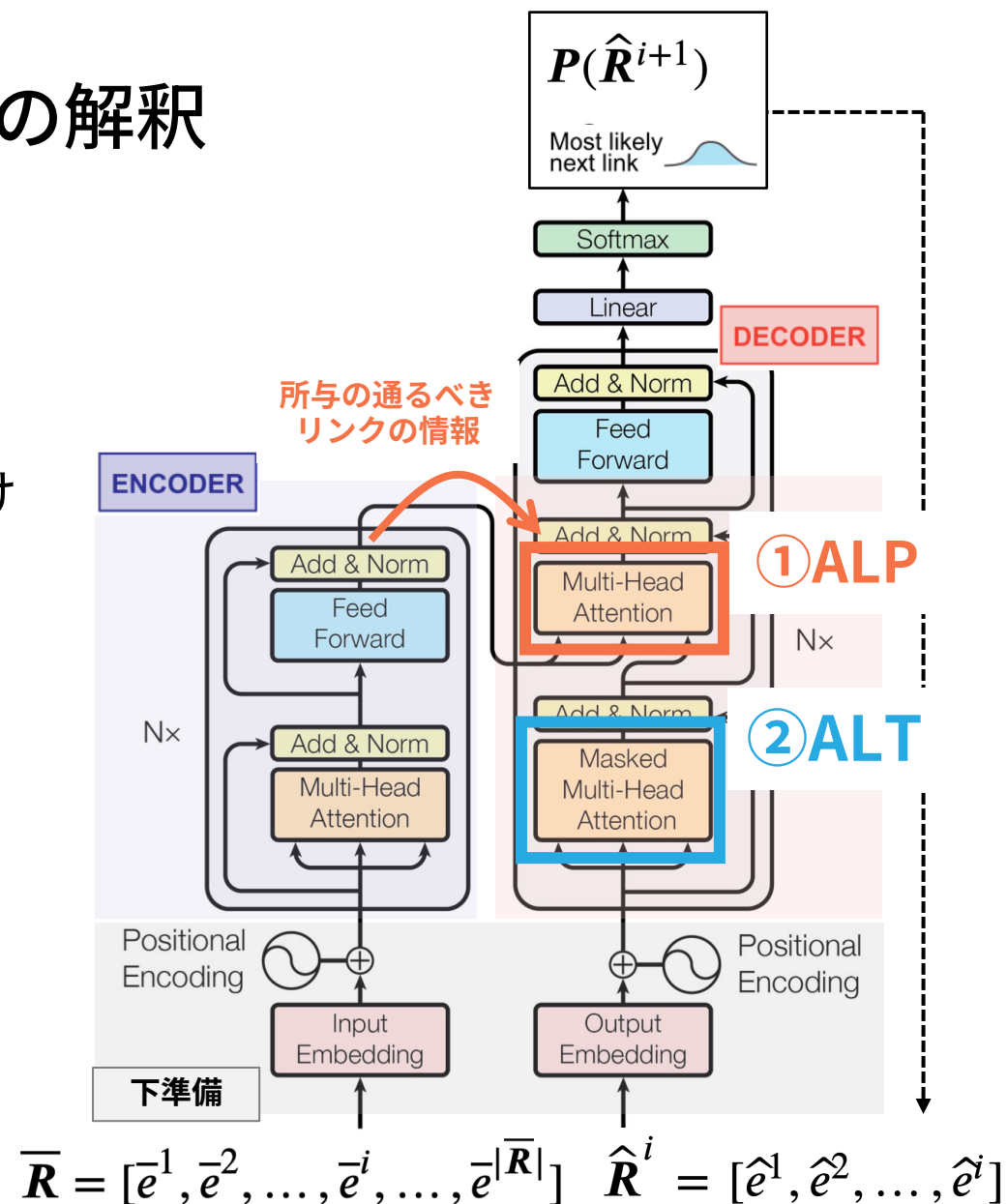
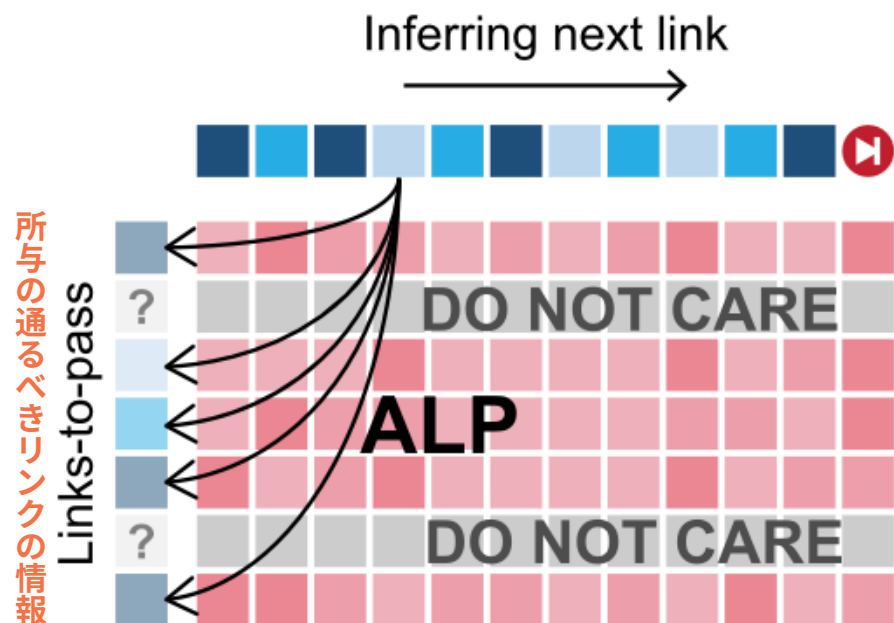
4.2 エンコーダー・デコーダー



4.3 RoutesFormerにおけるAttentionの解釈

① Attention to links-to-pass (ALP)

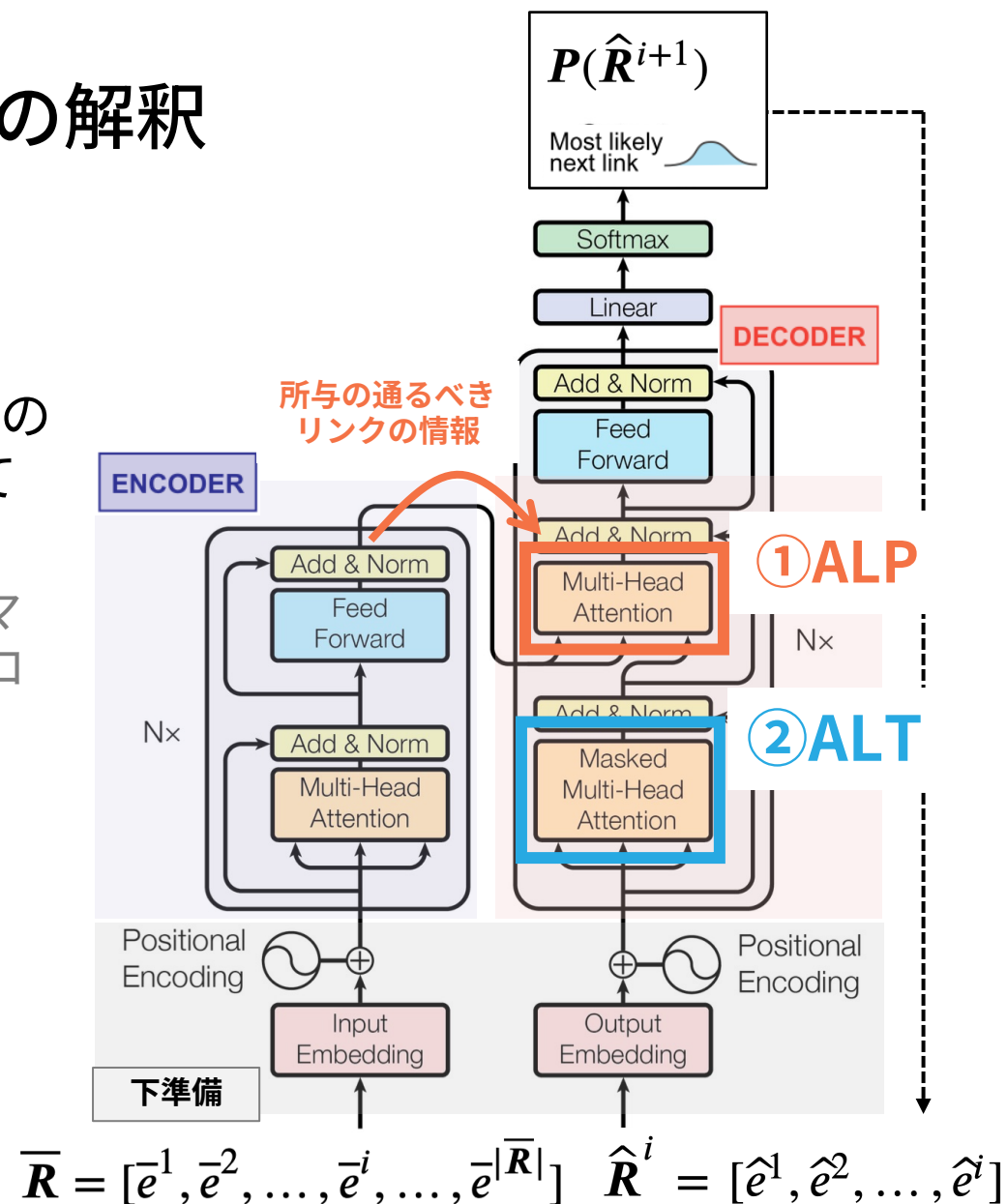
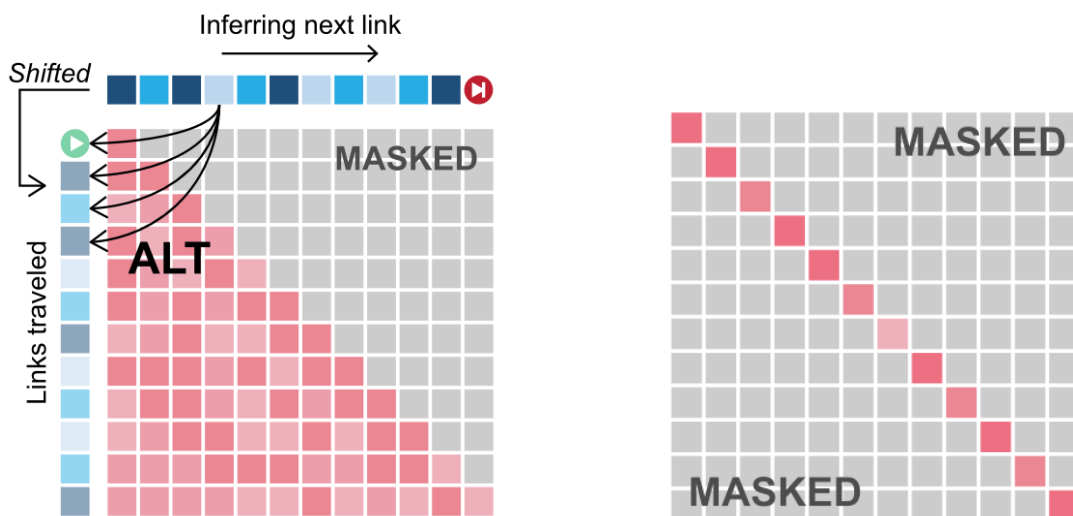
- Decoder内のCross-Attention
- $i+1$ 番目のリンクの推論時、スタートからゴールまでの**所与の通るべきリンク**に、どの程度注意を向けているか（＝そこからどの程度情報を得ているか）



4.3 RoutesFormerにおけるAttentionの解釈

② Attention to links traveled (ALT)

- Decoder内のSelf-Attention
- $i+1$ 番目のリンクの推論時、スタートから i 番目までのこれまで辿ってきたリンクにどの程度注意を向けているか（＝そこからどの程度情報を得ているか）
- $i+1$ 番目のリンク推論時に i 番目のリンク以外への注意をマスクすると、現在のリンクのみから情報を得る（＝マルコフ仮定）状況になる（右下図）→5.5節 Ablation Study



4.4 補足事項

- 次のリンクの選択確率を出した後、**隣接するリンクの中で最も確率の高いものを出力する**
- 学習時：クロスエントロピー誤差を用いてパラメタ調整
- 推論時：学習したパラメータで推論
 - **所与のリンクを通らない結果となる可能性もある**
→その場合は所与リンクペア間の部分問題に分割して解く

if $\bar{R}$ is not the subsequence of $\hat{R}$ then

$$\left| \hat{R} \leftarrow \text{RFI}(G, [\bar{e}^1, \bar{e}^2], \hat{\theta}, \theta_{\text{hyper}}) \oplus \cdots \oplus \text{RFI}(G, [\bar{e}^{|\bar{R}|-1}, \bar{e}^{|\bar{R}|}], \hat{\theta}, \theta_{\text{hyper}}) \right.$$

言語モデルをそのまま適応しているため、「所与のリンクを必ず通る」「隣接したリンクを選ぶ」という経路選択特有の制約条件にモデルの中からアプローチすることはできていない

5. Case study

5.1 データ設定・比較対象

- 2011年3月の、上海のタクシー8407台のGPSデータ
 - GPSのサンプリング間隔10秒
 - 人を乗せている区間のデータのみ用いる
- 80%を訓練用、20%をテスト用データとして用いる
- 疎な軌跡データの再現：(a)と(b)の混合
 - (a) 10%～90%の割合でデータをマスクする
 - (b) **GPSのサンプリング間隔が【60秒・180秒・300秒・420秒・0とDのみ】**である状況を再現してマスクする

比較対象のモデル

- Path size logit (PSL)
- DNN-PSL
- Recursive Logit (RL)
- Adversarial Inverse Reinforcement Learning (AIRL)
- Generative Adversarial Imitation Learning (GAIL) and Behavioral Cloning (BC)
- Dynamic LSTM (D-LSTM)

5.2 評価指標

- Total Link Length Accuracy (TLLA)
【(推論パスのリンク長)/(正解パスのリンク長)】
の平均

$$TLLA = \frac{1}{|P_{\text{test}}|} \sum_{n=1}^{|P_{\text{test}}|} \frac{\sum_{e \in \hat{R}_n \cap R_n} e.\text{length}}{\sum_{e \in R_n} e.\text{length}}$$

- Edit Distance (ED)
【(add,delete,modifyの必要回数)/(リンク数)】
の平均

$$ED = \frac{1}{|P_{\text{test}}|} \sum_{n=1}^{|P_{\text{test}}|} \min\left(\frac{\text{Edit}(\hat{R}_n, R_n)}{|R_n|}, 1\right)$$

- BiLingual Evaluation Understudy score (BLEU)
 - N-gramのマッチング度
- Jensen-Shannon Divergence
 - 確率分布の類似性

$$BLEU = \frac{1}{|P_{\text{test}}|} \sum_{n=1}^{|P_{\text{test}}|} \text{Bleu}(\hat{R}_n, R_n)$$

$$\text{JSD}(P, Q) = \frac{1}{2} \text{KLD}\left(P, \frac{P+Q}{2}\right) + \frac{1}{2} \text{KLD}\left(Q, \frac{P+Q}{2}\right)$$

$$\text{KLD}(P, Q) = \sum_n P_n \log \frac{P_n}{Q_n}$$

- Path completion rate (PCR)
 - 【所与リンクを繋げるのに成功したパス数/全パス数】

5.3 推論精度の比較

- RoutesFormerは**全てのタスク・全ての精度指標で旧来のモデルを凌駕**

- PCR=所与のリンクを通れるか は最下位

- 他のモデルは、タスクにより順位が入れ替わる

DNN-PSL > PSL

- DNNで非線形効用をうまく近似
- 間隔60sでは逆転：短いパスでの効用の違い

PSL > RL

- 実情に反した多くの経路の考慮

AIRL :

- サンプルの非効率性、いい関数を見つけづらい

Intervals	Models	Single-destination					Multi-destination				
		TLLA↑	ED↓	BLEU↑	JSD↓	PCR↑	TLLA↑	ED↓	BLEU↑	JSD↓	PCR↑
60 s	PSL (1999)	0.833 ^{#6}	0.195 ^{#6}	0.826 ^{#6}	0.301 ^{#6}	1.000 ^{#1}	0.879 ^{#7}	0.142 ^{#7}	0.888 ^{#7}	0.391 ^{#6}	1.000 ^{#1}
	DNN-PSL (2020)	0.812 ^{#7}	0.232 ^{#7}	0.811 ^{#7}	0.332 ^{#7}	1.000 ^{#1}	0.882 ^{#6}	0.136 ^{#6}	0.889 ^{#6}	0.448 ^{#7}	1.000 ^{#1}
	RL (2013)	0.901 ^{#5}	0.100 ^{#5}	0.898 ^{#5}	0.215 ^{#5}	1.000 ^{#1}	0.967 ^{#2}	0.047 ^{#3}	0.963 ^{#2}	0.109 ^{#2}	1.000 ^{#1}
	BC (2022)	0.921 ^{#4}	0.080 ^{#4}	0.924 ^{#4}	0.213 ^{#4}	0.851 ^{#7}	0.956 ^{#4}	0.062 ^{#5}	0.945 ^{#4}	0.146 ^{#5}	0.792 ^{#7}
	AIRL (2022)	0.976 ^{#2}	0.029 ^{#2}	0.972 ^{#2}	0.081 ^{#2}	0.995 ^{#4}	0.960 ^{#3}	0.046 ^{#2}	0.958 ^{#3}	0.121 ^{#3}	0.988 ^{#4}
	D-LSTM (2022)	0.961 ^{#3}	0.054 ^{#3}	0.950 ^{#3}	0.138 ^{#3}	0.981 ^{#5}	0.944 ^{#5}	0.050 ^{#4}	0.941 ^{#5}	0.133 ^{#4}	0.985 ^{#5}
	RoutesFormer	0.982^{#1}	0.028^{#1}	0.975^{#1}	0.074^{#1}	0.947^{#6}	0.975^{#1}	0.037^{#1}	0.969^{#1}	0.101^{#1}	0.951^{#6}
180 s	PSL (1999)	0.757 ^{#7}	0.234 ^{#6}	0.778 ^{#6}	0.599 ^{#6}	1.000 ^{#1}	0.840 ^{#6}	0.175 ^{#7}	0.833 ^{#7}	0.482 ^{#7}	1.000 ^{#1}
	DNN-PSL (2020)	0.820 ^{#5}	0.171 ^{#4}	0.839 ^{#4}	0.413 ^{#4}	1.000 ^{#1}	0.849 ^{#5}	0.171 ^{#6}	0.838 ^{#6}	0.440 ^{#6}	1.000 ^{#1}
	RL (2013)	0.626 ^{#6}	0.389 ^{#7}	0.577 ^{#7}	0.760 ^{#7}	1.000 ^{#1}	0.821 ^{#7}	0.157 ^{#4}	0.833 ^{#5}	0.336 ^{#3}	1.000 ^{#1}
	BC (2022)	0.824 ^{#4}	0.184 ^{#5}	0.830 ^{#5}	0.506 ^{#5}	0.877 ^{#7}	0.858 ^{#4}	0.165 ^{#5}	0.850 ^{#3}	0.369 ^{#5}	0.843 ^{#7}
	AIRL (2022)	0.910 ^{#3}	0.090 ^{#1}	0.911 ^{#2}	0.248 ^{#2}	0.993 ^{#4}	0.863 ^{#3}	0.138 ^{#2}	0.871 ^{#2}	0.318 ^{#2}	0.990 ^{#4}
	D-LSTM (2022)	0.911 ^{#2}	0.123 ^{#3}	0.885 ^{#3}	0.252 ^{#3}	0.986 ^{#5}	0.887 ^{#2}	0.146 ^{#3}	0.843 ^{#4}	0.355 ^{#4}	0.988 ^{#5}
	RoutesFormer	0.922^{#1}	0.090^{#1}	0.916^{#1}	0.218^{#1}	0.961^{#6}	0.906^{#1}	0.105^{#1}	0.906^{#1}	0.246^{#1}	0.966^{#6}
300 s	PSL (1999)	0.770 ^{#6}	0.243 ^{#6}	0.768 ^{#6}	0.558 ^{#5}	1.000 ^{#1}	0.780 ^{#6}	0.222 ^{#5}	0.781 ^{#5}	0.479 ^{#7}	1.000 ^{#1}
	DNN-PSL (2020)	0.835 ^{#4}	0.173 ^{#4}	0.831 ^{#4}	0.339 ^{#4}	1.000 ^{#1}	0.791 ^{#5}	0.217 ^{#4}	0.787 ^{#4}	0.468 ^{#5}	1.000 ^{#1}
	RL (2013)	0.548 ^{#7}	0.440 ^{#7}	0.520 ^{#7}	0.876 ^{#7}	1.000 ^{#1}	0.765 ^{#7}	0.225 ^{#6}	0.771 ^{#7}	0.439 ^{#4}	1.000 ^{#1}
	BC (2022)	0.802 ^{#5}	0.222 ^{#5}	0.795 ^{#5}	0.598 ^{#6}	0.888 ^{#7}	0.799 ^{#4}	0.239 ^{#7}	0.780 ^{#6}	0.477 ^{#6}	0.867 ^{#7}
	AIRL (2022)	0.877 ^{#2}	0.133 ^{#2}	0.868 ^{#2}	0.318 ^{#3}	0.990 ^{#4}	0.811 ^{#3}	0.209 ^{#3}	0.802 ^{#3}	0.420 ^{#3}	0.993 ^{#4}
	D-LSTM (2022)	0.872 ^{#3}	0.141 ^{#3}	0.868 ^{#2}	0.290 ^{#2}	0.982 ^{#5}	0.831 ^{#2}	0.181 ^{#2}	0.833 ^{#2}	0.397 ^{#2}	0.989 ^{#5}
	RoutesFormer	0.902^{#1}	0.113^{#1}	0.895^{#1}	0.258^{#1}	0.970^{#6}	0.860^{#1}	0.158^{#1}	0.857^{#1}	0.325^{#1}	0.972^{#6}
420 s	PSL (1999)	0.736 ^{#6}	0.251 ^{#6}	0.760 ^{#6}	0.546 ^{#5}	1.000 ^{#1}	0.755 ^{#5}	0.250 ^{#4}	0.752 ^{#5}	0.494 ^{#6}	1.000 ^{#1}
	DNN-PSL (2020)	0.798 ^{#4}	0.180 ^{#4}	0.823 ^{#4}	0.305 ^{#3}	1.000 ^{#1}	0.760 ^{#4}	0.244 ^{#3}	0.758 ^{#4}	0.481 ^{#5}	1.000 ^{#1}
	RL (2013)	0.501 ^{#7}	0.460 ^{#7}	0.497 ^{#7}	0.887 ^{#7}	1.000 ^{#1}	0.743 ^{#7}	0.261 ^{#6}	0.736 ^{#7}	0.480 ^{#4}	1.000 ^{#1}
	BC (2022)	0.777 ^{#5}	0.244 ^{#5}	0.772 ^{#5}	0.619 ^{#6}	0.889 ^{#7}	0.747 ^{#6}	0.280 ^{#7}	0.739 ^{#6}	0.521 ^{#7}	0.869 ^{#7}
	AIRL (2022)	0.851 ^{#3}	0.163 ^{#3}	0.838 ^{#3}	0.335 ^{#4}	0.995 ^{#4}	0.765 ^{#3}	0.250 ^{#4}	0.762 ^{#3}	0.463 ^{#3}	0.992 ^{#4}
	D-LSTM (2022)	0.862 ^{#2}	0.129 ^{#2}	0.860 ^{#2}	0.286 ^{#2}	0.989 ^{#5}	0.794 ^{#2}	0.223 ^{#2}	0.795 ^{#2}	0.407 ^{#2}	0.986 ^{#5}
	RoutesFormer	0.890^{#1}	0.125^{#1}	0.884^{#1}	0.268^{#1}	0.977^{#6}	0.819^{#1}	0.190^{#1}	0.828^{#1}	0.360^{#1}	0.973^{#6}
Only OD	PSL (1999)	0.790 ^{#5}	0.230 ^{#5}	0.789 ^{#5}	0.616 ^{#5}	1.000 ^{#1}	0.755 ^{#4}	0.251 ^{#4}	0.753 ^{#4}	0.496 ^{#4}	1.000 ^{#1}
	DNN-PSL (2020)	0.803 ^{#4}	0.221 ^{#4}	0.792 ^{#4}	0.362 ^{#3}	1.000 ^{#1}	0.756 ^{#3}	0.244 ^{#3}	0.760 ^{#3}	0.484 ^{#3}	1.000 ^{#1}
	RL (2013)	0.459 ^{#7}	0.489 ^{#7}	0.466 ^{#7}	0.935 ^{#7}	1.000 ^{#1}	0.670 ^{#6}	0.327 ^{#6}	0.668 ^{#7}	0.567 ^{#6}	1.000 ^{#1}
	BC (2022)	0.711 ^{#6}	0.274 ^{#6}	0.746 ^{#6}	0.638 ^{#6}	0.989 ^{#6}	0.654 ^{#7}	0.340 ^{#7}	0.680 ^{#6}	0.584 ^{#7}	0.982 ^{#7}
	AIRL (2022)	0.806 ^{#3}	0.191 ^{#3}	0.811 ^{#3}	0.367 ^{#4}	1.000 ^{#1}	0.755 ^{#4}	0.307 ^{#5}	0.734 ^{#5}	0.524 ^{#5}	0.998 ^{#4}
	D-LSTM (2022)	0.873 ^{#2}	0.141 ^{#2}	0.870 ^{#2}	0.336 ^{#2}	0.994 ^{#5}	0.770 ^{#2}	0.240 ^{#2}	0.762 ^{#2}	0.451 ^{#2}	0.996 ^{#5}
	RoutesFormer	0.876^{#1}	0.140^{#1}	0.874^{#1}	0.329^{#1}	0.989^{#6}	0.791^{#1}	0.234^{#1}	0.786^{#1}	0.436^{#1}	0.993^{#6}

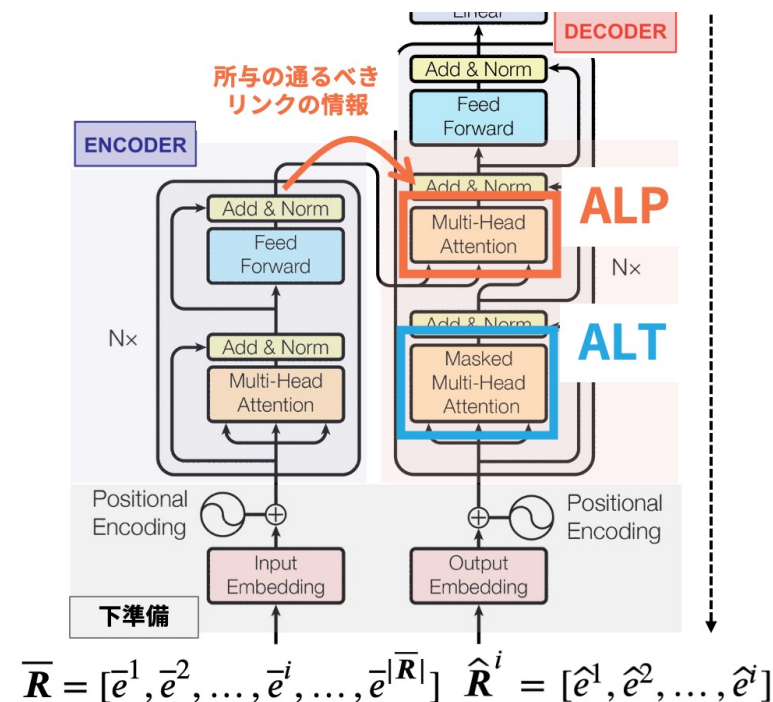
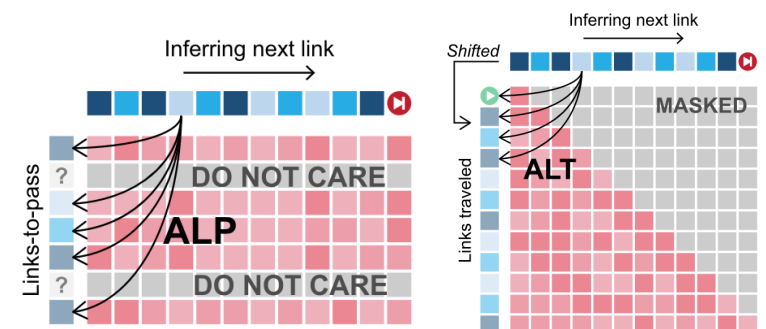
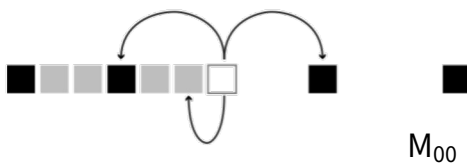
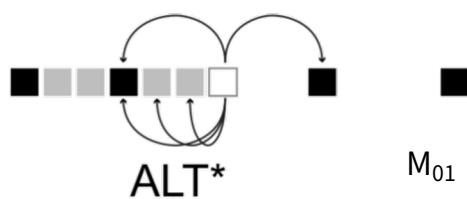
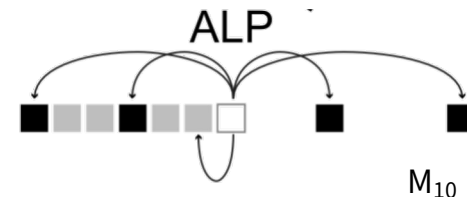
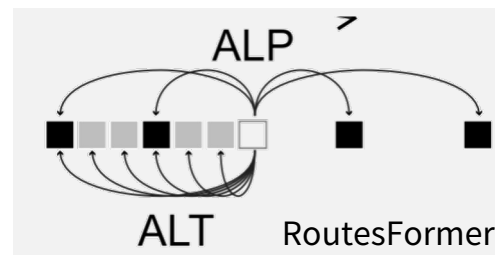
5.4 計算時間の比較

- 並列処理可能なRoutesFormerは、**パラメタ数の多さの割に短い学習時間**
 - ⇔ **推論時間は長い**
- パスベースのPSL, DNN-PSLはメモリを使ってパス候補を事前に格納しているため、最短の学習・推論時間
- 古典的なRLはモデルサイズ・計算時間とも小さい
- AIRLは逆強化学習という特性上広範なサンプル収集が必要 → サイズの割に長い学習時間

Models	Parameters	Training time(s) under different training size				Inference time(s)
		1k	2k	4k	8k	
PSL (1999)	13	30.5	52.4	81.1	170.3	0.00185
DNN-PSL (2020)	5249	41.7	86.3	157.7	289.4	0.00247
RL (2013)	56	417.2	1121.8	1910.3	2912.1	0.0072
BC (2022)	41 996	709.5	1346.2	2478.5	5021.3	0.0352
AIRL (2022)	41996(actor)	9935.6	13 948.1	15 822.9	19 204.2	0.0381
D-LSTM (2022)	803 629	945.7	2023.1	3981.6	8281.0	0.0498
RoutesFormer	2 260 653	2858.2	5943.0	10 531.1	19 140.5	0.436

5.5 Ablation study

- ALP, ALTができないようにした**切除モデル**で検証し、**Attention機構の意義**を分析する
 - M_{10} : ALPはできるが、ALTでは**現在リンクからしか情報を得られない**
 - M_{01} : ALPは**推論中の前後の所与リンクのみ**に向けられ、**ALTも最後に通った所与リンクまで**からしか情報を得られない
 - M_{00} : ALPは**推論中の前後の所与リンクのみ**に向けられ、ALTでは**現在リンクからしか情報を得られない**



5.5.1 精度の比較

#で書かれた上付き文字は全モデルとの比較における順位

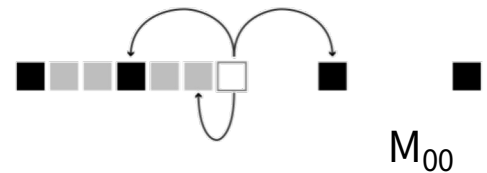
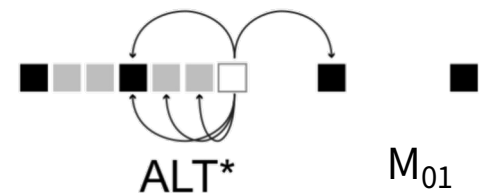
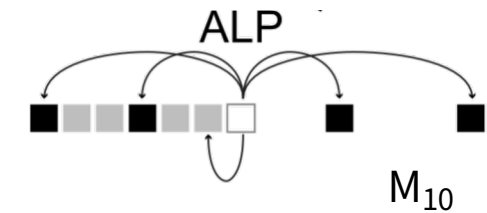
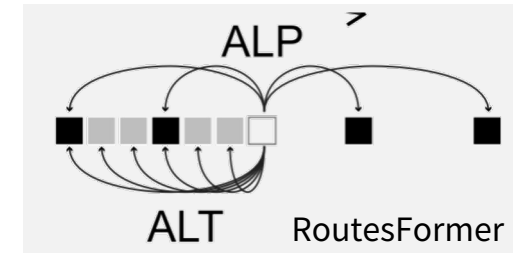
Sampling intervals /Sampled links	Models	Single-destination				Multi-destination			
		TLLA↑	ED↓	BLEU↑	JSD↓	TLLA↑	ED↓	BLEU↑	JSD↓
60 s	M_{00}	0.946 ^{#4}	0.078 ^{#4}	0.931 ^{#4}	0.195 ^{#4}	0.934 ^{#6}	0.081 ^{#6}	0.927 ^{#6}	0.172 ^{#6}
	M_{01}	0.964 ^{#3}	0.038 ^{#3}	0.964 ^{#3}	0.096 ^{#3}	0.950 ^{#4}	0.041 ^{#2}	0.945 ^{#4}	0.124 ^{#4}
	M_{10}	0.968 ^{#3}	0.066 ^{#4}	0.942 ^{#4}	0.180 ^{#4}	0.955 ^{#5}	0.067 ^{#6}	0.949 ^{#4}	0.153 ^{#6}
	M_{11}	0.982^{#1}	0.028^{#1}	0.975^{#1}	0.074^{#1}	0.975^{#1}	0.037^{#1}	0.969^{#1}	0.101^{#1}
180 s	M_{00}	0.889 ^{#4}	0.161 ^{#4}	0.842 ^{#4}	0.369 ^{#4}	0.868 ^{#3}	0.157 ^{#4}	0.861 ^{#3}	0.383 ^{#6}
	M_{01}	0.910 ^{#3}	0.109 ^{#3}	0.903 ^{#3}	0.256 ^{#4}	0.887 ^{#2}	0.130 ^{#2}	0.884 ^{#2}	0.375 ^{#5}
	M_{10}	0.894 ^{#4}	0.159 ^{#4}	0.858 ^{#4}	0.351 ^{#4}	0.872 ^{#3}	0.145 ^{#3}	0.878 ^{#3}	0.314 ^{#2}
	M_{11}	0.922^{#1}	0.090^{#1}	0.916^{#1}	0.218^{#1}	0.906^{#1}	0.105^{#1}	0.906^{#1}	0.246^{#1}
300 s	M_{00}	0.862 ^{#4}	0.203 ^{#5}	0.811 ^{#5}	0.421 ^{#5}	0.826 ^{#3}	0.241 ^{#7}	0.816 ^{#3}	0.411 ^{#3}
	M_{01}	0.890 ^{#2}	0.124 ^{#2}	0.886 ^{#2}	0.281 ^{#2}	0.842 ^{#2}	0.172 ^{#2}	0.835 ^{#2}	0.386 ^{#2}
	M_{10}	0.871 ^{#4}	0.189 ^{#4}	0.833 ^{#4}	0.412 ^{#4}	0.813 ^{#3}	0.208 ^{#3}	0.821 ^{#3}	0.403 ^{#3}
	M_{11}	0.902^{#1}	0.113^{#1}	0.895^{#1}	0.258^{#1}	0.860^{#1}	0.158^{#1}	0.857^{#1}	0.325^{#1}
420 s	M_{00}	0.843 ^{#4}	0.226 ^{#5}	0.796 ^{#5}	0.465 ^{#5}	0.742 ^{#8}	0.264 ^{#7}	0.751 ^{#6}	0.487 ^{#5}
	M_{01}	0.870 ^{#2}	0.136 ^{#3}	0.877 ^{#2}	0.308 ^{#4}	0.816 ^{#2}	0.202 ^{#2}	0.824 ^{#2}	0.390 ^{#2}
	M_{10}	0.859 ^{#3}	0.204 ^{#5}	0.820 ^{#5}	0.425 ^{#5}	0.811 ^{#2}	0.245 ^{#3}	0.788 ^{#3}	0.443 ^{#3}
	M_{11}	0.890^{#1}	0.125^{#1}	0.884^{#1}	0.268^{#1}	0.819^{#1}	0.190^{#1}	0.828^{#1}	0.360^{#1}
OD only	M_{00}	0.836 ^{#3}	0.253 ^{#7}	0.781 ^{#6}	0.531 ^{#5}	0.753 ^{#5}	0.298 ^{#5}	0.740 ^{#5}	0.526 ^{#6}
	M_{01}	0.876 ^{#1}	0.140 ^{#1}	0.874 ^{#1}	0.329 ^{#1}	0.791 ^{#1}	0.234 ^{#1}	0.786 ^{#1}	0.436 ^{#1}
	M_{10}	0.836 ^{#3}	0.253 ^{#7}	0.781 ^{#6}	0.531 ^{#5}	0.753 ^{#5}	0.298 ^{#5}	0.740 ^{#5}	0.526 ^{#6}
	M_{11}	0.876^{#1}	0.140^{#1}	0.874^{#1}	0.329^{#1}	0.791^{#1}	0.234^{#1}	0.786^{#1}	0.436^{#1}

5.5.1 精度の比較

- M_{01} が比較的高い精度 $\Leftrightarrow M_{10}$ 、 M_{00} の精度は低い
 - → **ALTの有無がより大きく影響**
 - 単一目的地のデータでは特に、過去の経路の影響大（ $\Leftrightarrow$ 複数目的地では多様な経路選択）
 - 疎なサンプリング間隔ほど、ALTが顕著な影響（ $\Leftrightarrow$ 密なサンプリング間隔では相対的にALPの影響が大きく）

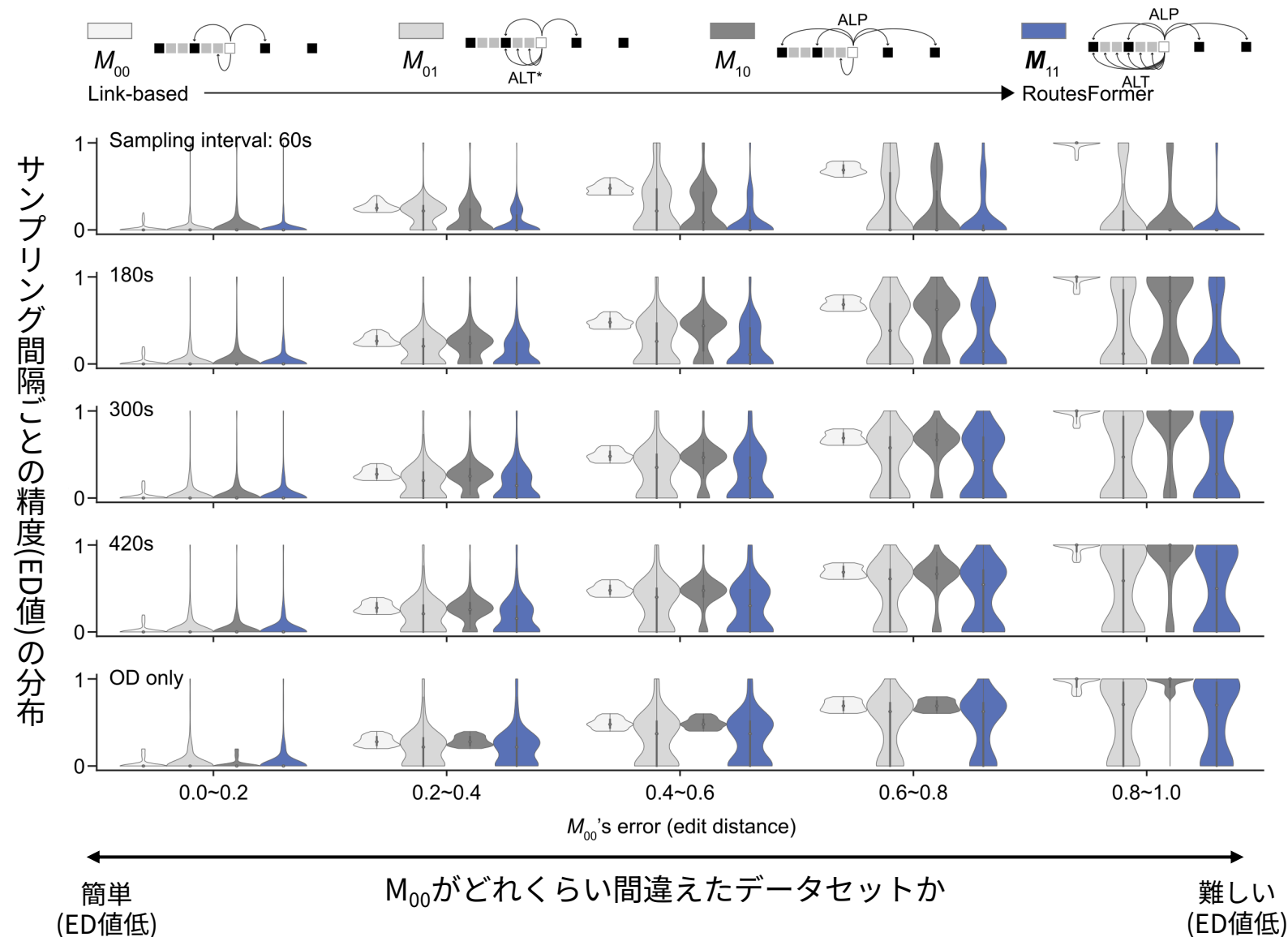
Edit Distanceの平均値（低いほど良い）

モデル	単一目的地データ	複数目的地データ
RoutesFormer	0.099	0.145
M_{10} (ALTなし)	0.174	0.193
M_{01} (ALPなし)	0.109	0.156
M_{00} (両方なし)	0.184	0.208



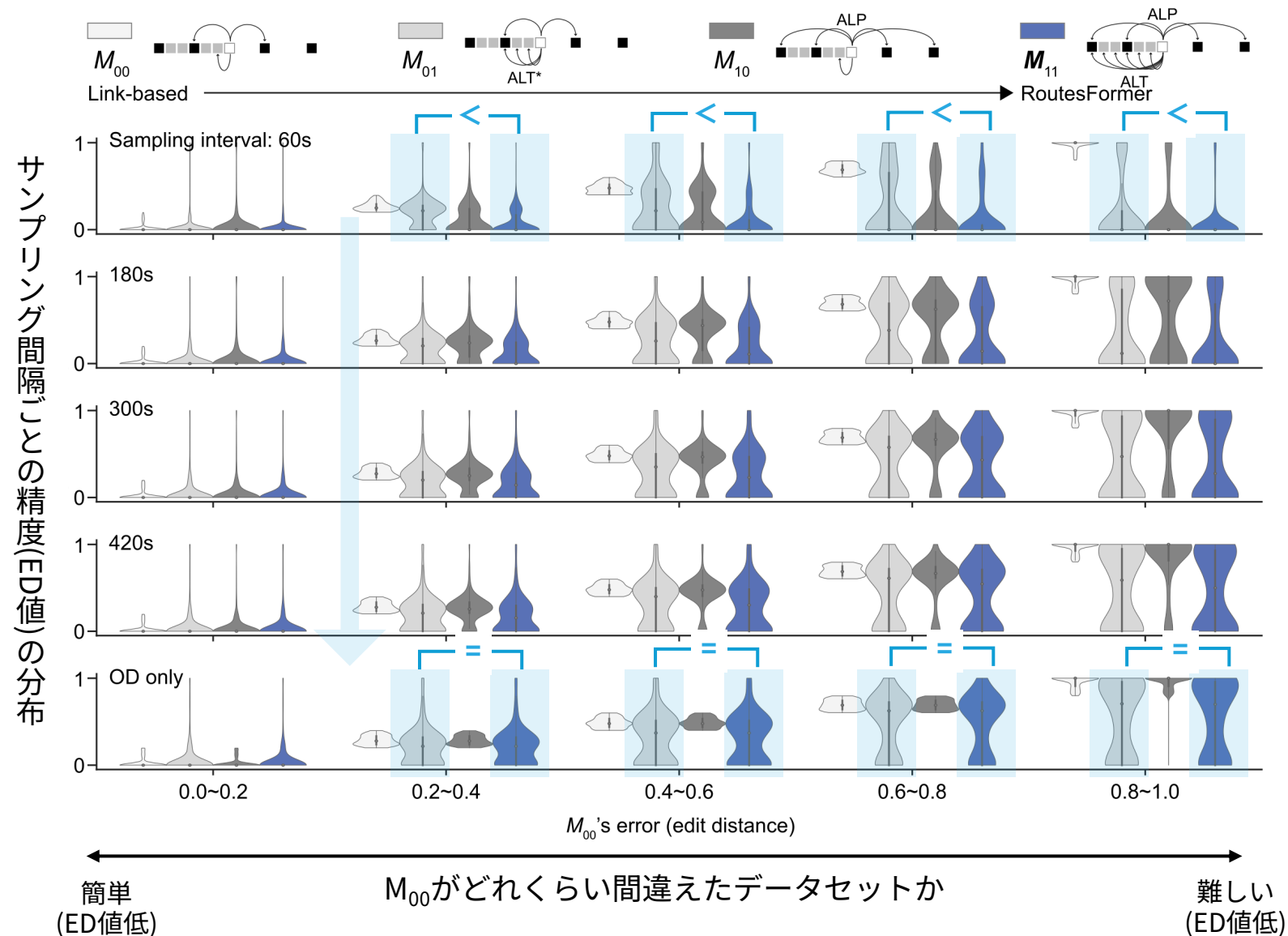
5.5.2 精度の分布比較

- Edit Distanceの分布を、 M_{00} の精度に基づいて区分したデータセットグループごとに、各モデルについて可視化
- 分布の**上が細く下が太い**（＝低いED値が多い）ほど良い



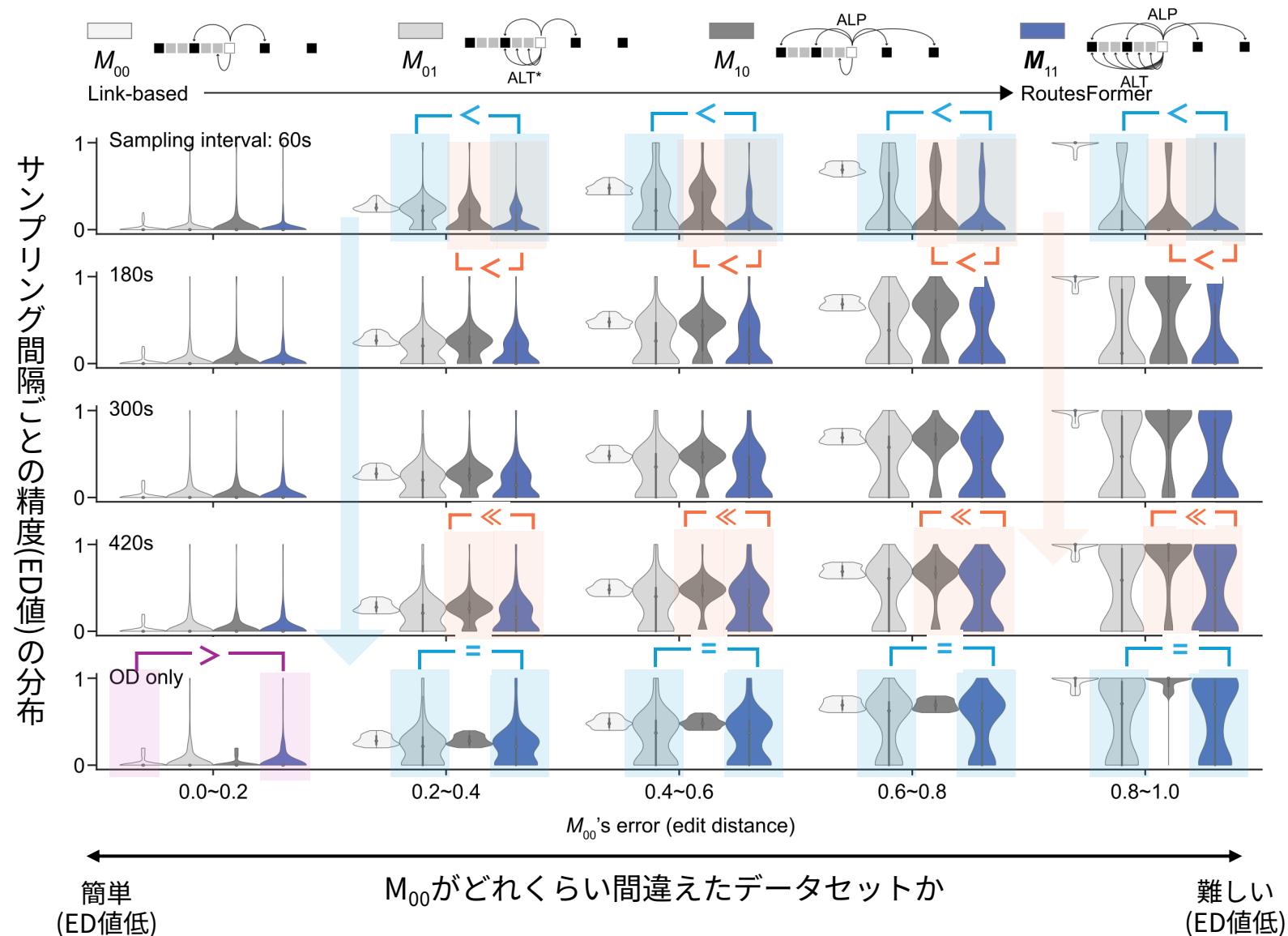
5.5.2 精度の分布比較

	M_{01} (ALPなし)
密なサンプル リング間隔	悪い
疎なサンプル リング間隔	本家に近づく



5.5.2 精度の分布比較

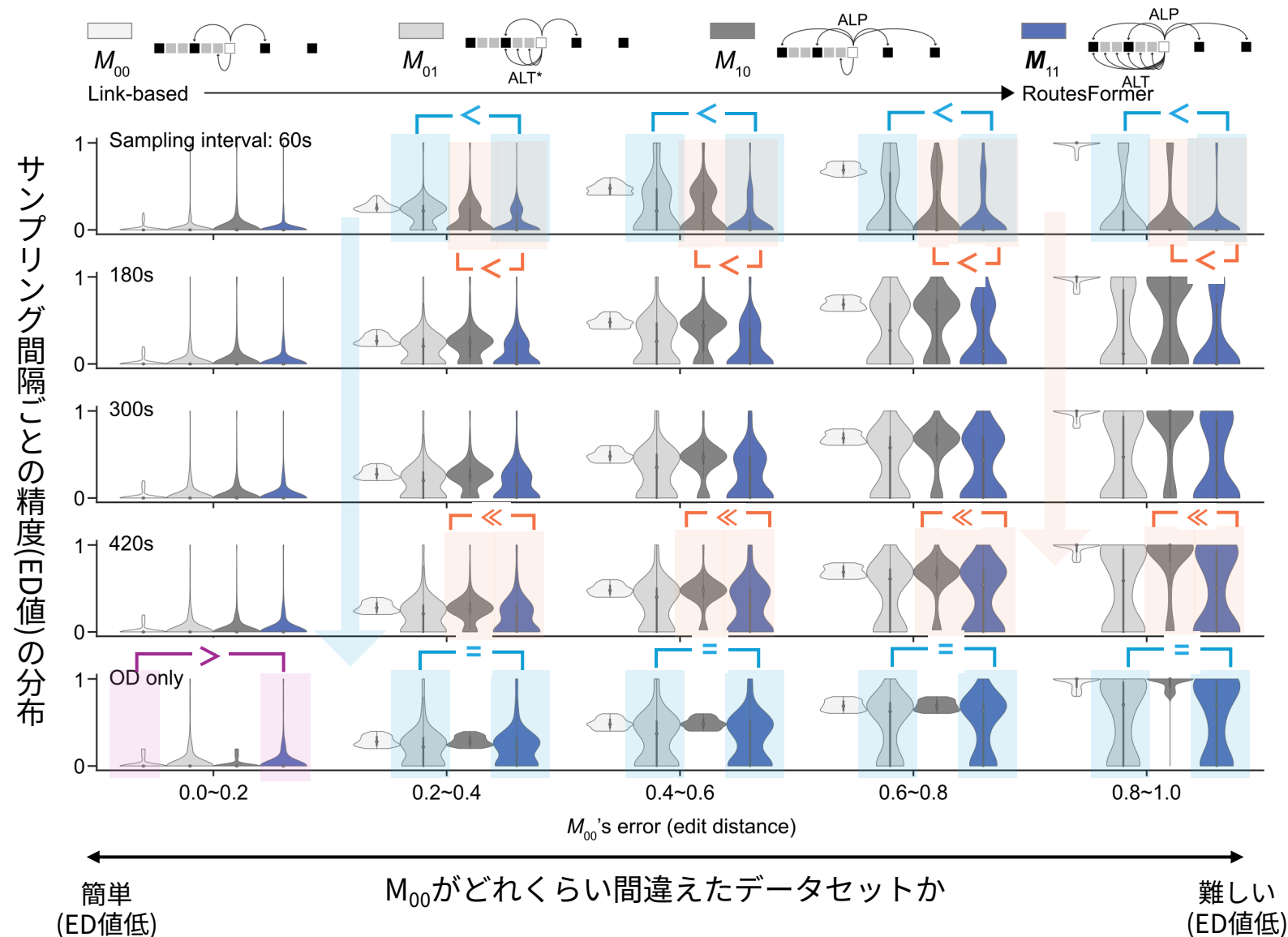
	M_{01} (ALPなし)	M_{10} (ALTなし)
密なサンプル リング間隔	悪い	悪い
疎なサンプル リング間隔	本家に近づく	より悪い



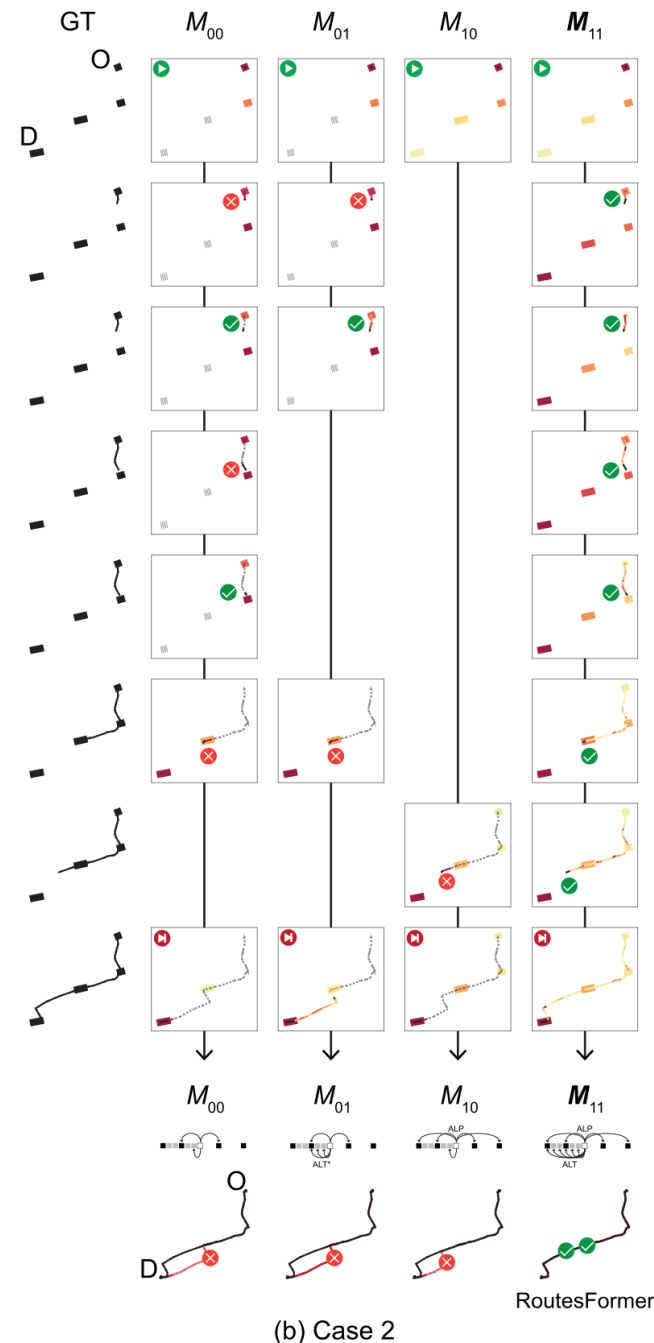
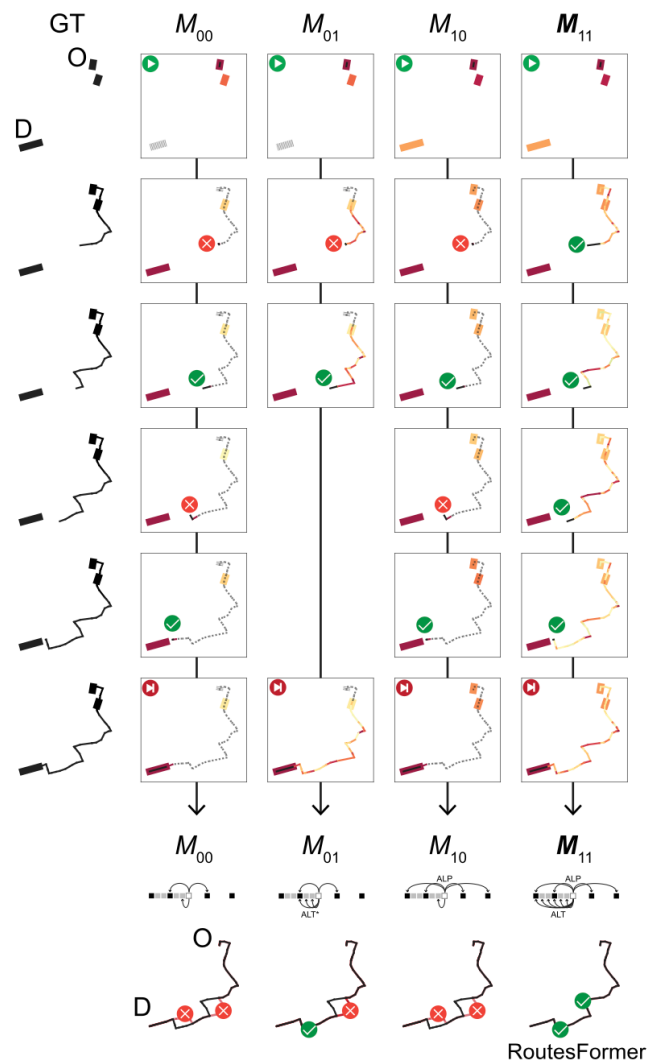
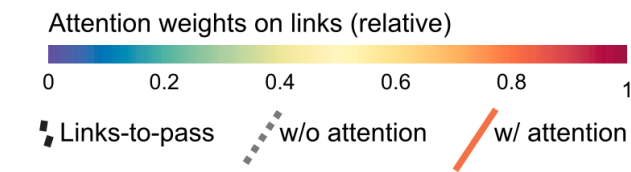
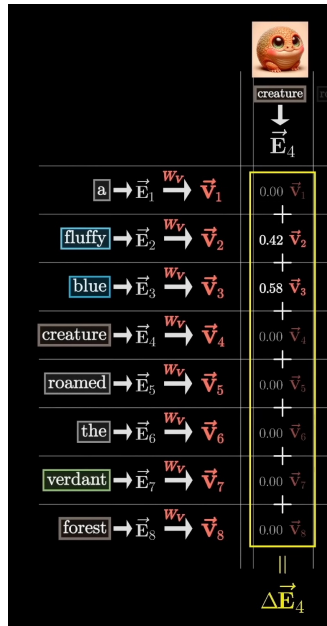
5.5.2 精度の分布比較

	M_{01} (ALPなし)	M_{10} (ALTなし)
密なサンプル リング間隔	悪い	悪い
疎なサンプル リング間隔	本家に近づく	より悪い

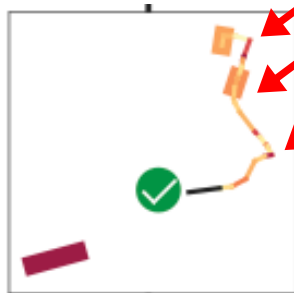
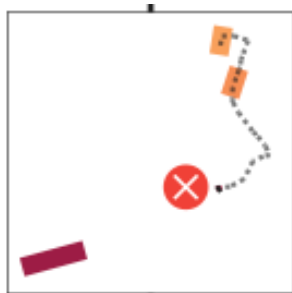
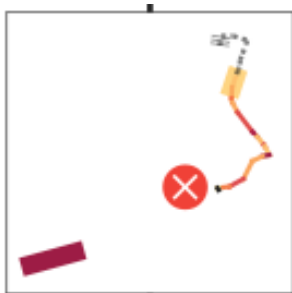
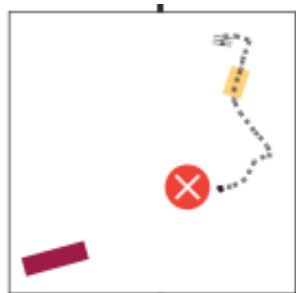
- 疎な間隔であるほど過去の情報を参照する余地が大きく、ALTの重要性が増す
- M_{00} の結果がいい場合も→嗜好の多様性



- あるリンクが推論されたとき、どのリンクの情報に特に注目して決定されたかは、各リンクへのattentionの値の大小を見ることでわかる
- 各リンクへのattentionの値（相対値）と、次のリンクの推論結果をステップごとに可視化
- 本家モデル以外は、「推論を誤ったステップ」と「その後正しい道に戻ったステップ」を抽出している

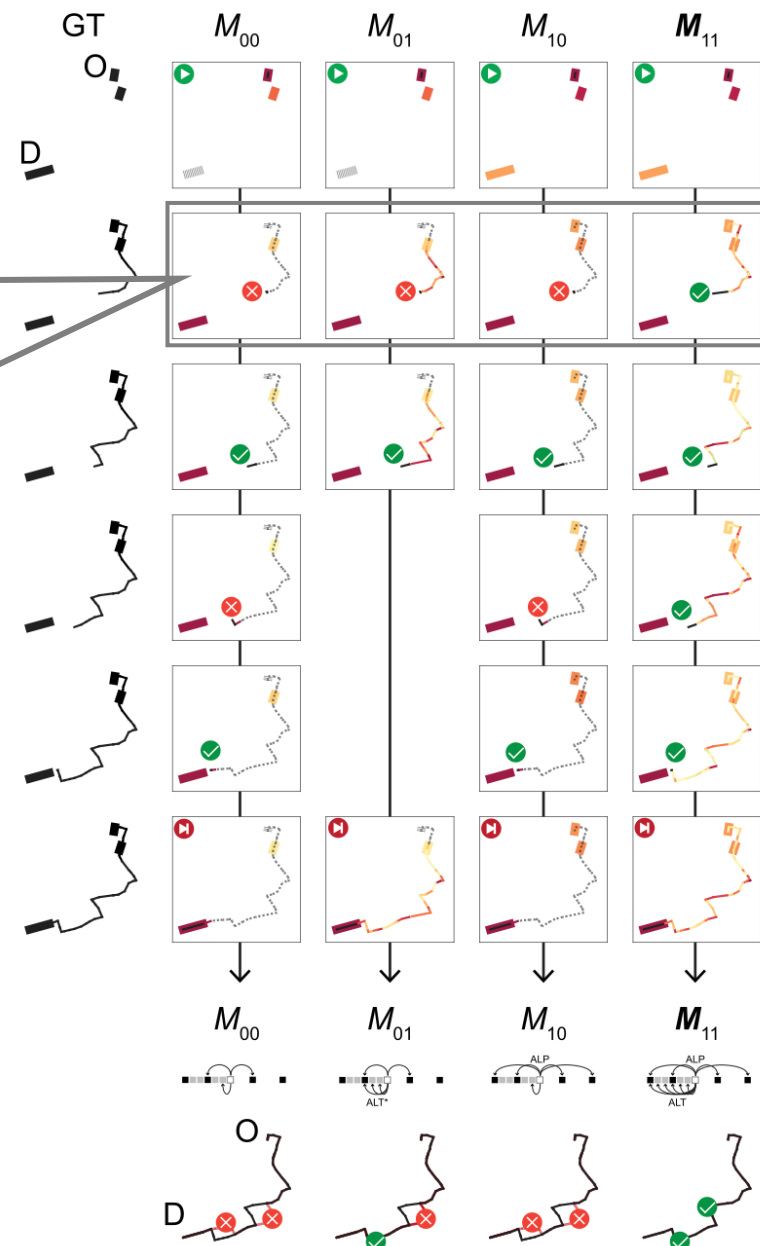


5.5.3 推論結果の比較

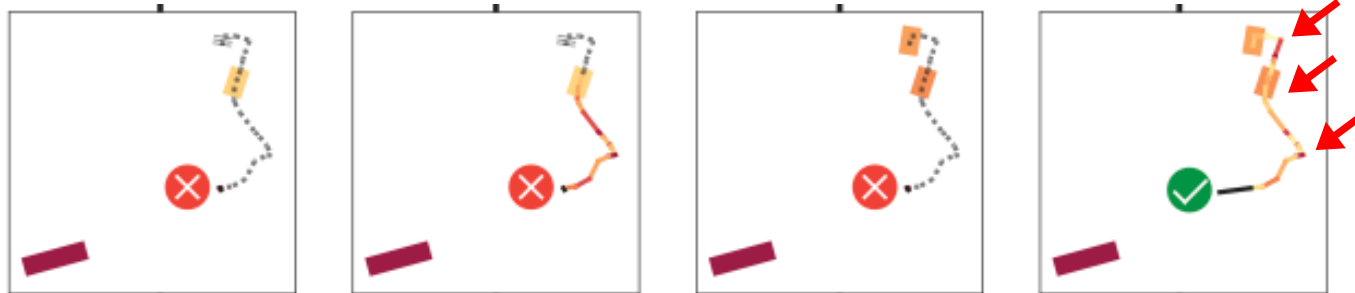


正解したRoutesFormerでは、現在リンクではなく、

- いくつかのより前に通ったリンク
- Oや、その次の通るべき所与(=太い)のリンクに、高いattentionを示している



5.5.3 推論結果の比較

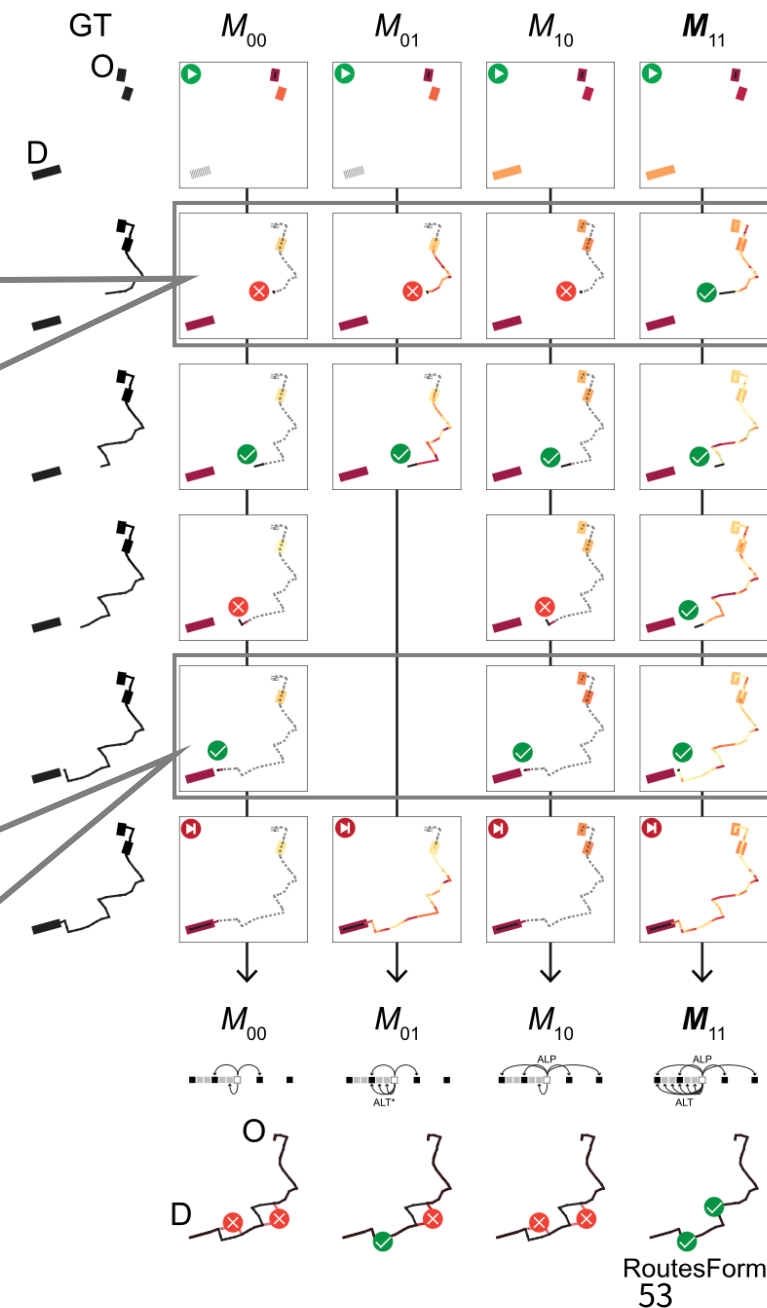


正解したRoutesFormerでは、現在リンクではなく、

- いくつかのより前に通ったリンク
- Oや、その次の通るべき所与(=太い)のリンクに、高いattentionを示している



ALTのないモデルで、局所的な回り道が発生
(このケース以外にも多数観察される)



6. Conclusion

6.1 Conclusion

- ✓ RoutesFormerは2つの面で既存の経路選択モデルの問題を解決した：
 - ①経路全体の情報を一括して扱える
 - ②マルコフ性を仮定しない
- ✓細かい特徴をモデル化でき、少ないデータ量でも**安定した学習**が可能
- ✓全タスクで**既存の経路選択モデルを上回る精度**を記録した
- ✓アブレーション分析により、**経路全体・過去の選択経路との関係を考慮する重要性を示す**とともに、現実世界の選択嗜好をよりよく理解する方法を提示した

6.2 発展可能性

今後のモデルの活用・改善：

- 既存の効用推定アーキテクチャよりも**言語モデルへのフィッティングに重き**を置いている
 - リンク特徴量は一般的なもののみで、ドメイン知識を入れていないので転移性が低い可能性
 - 全リンクを予測候補にした上で、隣接するものから選ぶという非合理性
- Attentionのメモリ消費と計算時間が大きい。①豊富で十分カバーされた訓練データがあり、②単純な仮定に収まらない多様な③連続的な選択の関係をみる向きに向いている

さらなる研究：

- ✓ 数少ない軌跡データから広大なネットワークの行動を学習する場合は、強化学習との組み合わせが考えられる
- ✓ ネットワーク間転移・都市間転移には、グラフ学習との組み合わせが有効
- ✓ より精緻な推論のため、**軌跡上の点の時空間情報の考慮**が重要

.....> ViRT (倉澤さん修論)

-