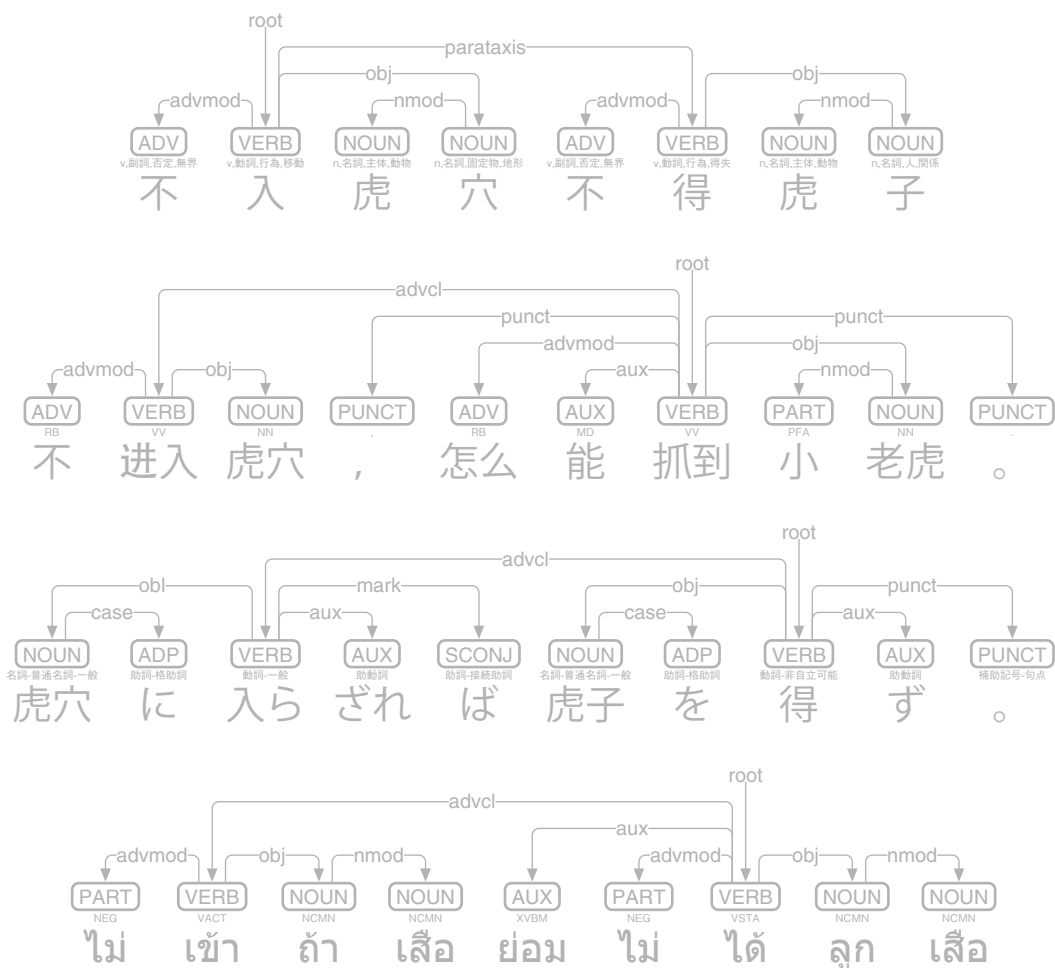


多言語情報処理

安岡孝一



2026年9月版

Universal Dependencies と BERT/RoBERTa/DeBERTa/GPT モデルによる多言語情報処理 (2026 年 9 月版)

安岡孝一 (京都大学人文科学研究所附属人文情報学創新センター)

1	概説編	5
1.1	Universal Dependencies と CoNLL-U	5
1.2	依存文法理論と Universal Dependencies	9
1.3	BERT/RoBERTa/DeBERTa モデルと穴埋めゲーム	14
1.4	GPT 系モデルと言語生成	16
1.5	Universal Dependencies における係り受け解析	17
2	古典中国語 (漢文) 編	21
2.1	古典中国語 UD における品詞付与と文切り	21
2.2	古典中国語 UD における係り受け解析	24
2.3	古典中国語ミニ DeBERTa モデルの製作	26
3	日本語編	30
3.1	国語研短単位にもとづく品詞付与と係り受け解析	30
3.2	国語研短単位ミニ DeBERTa モデルの製作	32
3.3	国語研長単位にもとづく品詞付与と係り受け解析	34
3.4	国語研長単位ミニ DeBERTa モデルの製作	36
3.5	文節にもとづく係り受け解析	37
4	英語編	38
4.1	英語 UD における品詞付与と係り受け解析	38
4.2	英語ミニ DeBERTa モデルの製作	39
4.3	英語 UD による構成素解析	41
5	フランス語編	43
5.1	フランス語 UD における品詞付与と係り受け解析	43
5.2	フランス語ミニ DeBERTa モデルの製作	44
6	ドイツ語編	47
6.1	ドイツ語 UD における品詞付与と係り受け解析	47
6.2	ドイツ語ミニ DeBERTa モデルの製作	48
7	オランダ語編	51
7.1	オランダ語 UD における品詞付与と係り受け解析	51
7.2	オランダ語ミニ DeBERTa モデルの製作	51
8	タイ語編	54
8.1	タイ語 UD における品詞付与	54
8.2	タイ語 UD における係り受け解析	55

8.3	タイ語ミニ DeBERTa モデルの製作	55
9	ベトナム語編	59
9.1	ベトナム語 UD における品詞付与	59
9.2	ベトナム語 UD における係り受け解析	60
9.3	ベトナム語ミニ DeBERTa モデルの製作	60
10	ロシア語編	63
10.1	ロシア語 UD における品詞付与と係り受け解析	63
10.2	ロシア語ミニ DeBERTa モデルの製作	63
11	ウクライナ語編	66
11.1	ウクライナ語 UD における品詞付与と係り受け解析	66
11.2	ウクライナ語ミニ DeBERTa モデルの製作	66
12	ベラルーシ語編	69
12.1	ベラルーシ語 UD における品詞付与と係り受け解析	69
12.2	ベラルーシ語ミニ DeBERTa モデルの製作	69
13	セルビア語編	72
13.1	セルビア語 UD における品詞付与と係り受け解析	72
13.2	セルビア語ミニ DeBERTa モデルの製作	73
14	クロアチア語編	76
14.1	クロアチア語 UD における品詞付与と係り受け解析	76
14.2	クロアチア語ミニ DeBERTa モデルの製作	76
15	コプト語編	79
15.1	コプト語 UD における品詞付与	79
15.2	コプト語 UD における係り受け解析	80
15.3	コプト語ミニ DeBERTa モデルの製作	80
16	スペイン語編	83
16.1	スペイン語 UD における品詞付与と係り受け解析	83
16.2	スペイン語ミニ DeBERTa モデルの製作	84
17	カタルーニャ語編	87
17.1	カタルーニャ語 UD における品詞付与と係り受け解析	87
17.2	カタルーニャ語ミニ DeBERTa モデルの製作	88
18	ガリシア語編	91
18.1	ガリシア語 UD における品詞付与と係り受け解析	91
18.2	ガリシア語ミニ DeBERTa モデルの製作	91

19 ポルトガル語編	94
19.1 ポルトガル語 UD における品詞付与と係り受け解析	94
19.2 ポルトガル語ミニ DeBERTa モデルの製作	94
20 イタリア語編	97
20.1 イタリア語 UD における品詞付与と係り受け解析	97
20.2 イタリア語ミニ DeBERTa モデルの製作	98
21 ラテン語編	100
21.1 ラテン語 UD における品詞付与と係り受け解析	100
21.2 ラテン語ミニ DeBERTa モデルの製作	100
22 バスク語編	103
22.1 バスク語 UD における品詞付与と係り受け解析	103
22.2 バスク語ミニ DeBERTa モデルの製作	104
23 トルコ語編	106
23.1 トルコ語 UD における品詞付与と係り受け解析	106
23.2 トルコ語ミニ DeBERTa モデルの製作	106
24 韓国語編	109
24.1 韓国語 UD における品詞付与と係り受け解析	109
24.2 韓国語の形態素にもとづく品詞付与と係り受け解析	110
24.3 韓国語ミニ DeBERTa モデルの製作	111
25 アイヌ語編	114
25.1 アイヌ語 UD における品詞付与と係り受け解析	114
25.2 アイヌ語ミニ DeBERTa モデルの製作	115

1 概説編

本書では、Universal Dependencies (UD)^[1]と呼ばれる多言語係り受けコーパス^[2]を用いて、自然言語処理のうち、特に、品詞付与と係り受け解析に関して解説します。これらの解析手法の背後には、現実にはややコシイ数式が横たわっているのですが、本書では、できる限り数式を使わずに、これらの解析手法を解説します。一方で、現代の自然言語処理は、その多くを BERT^[3]・RoBERTa^[4]・DeBERTa^[5]・ModernBERT^[6]などの事前学習モデルに拠っていることから、これら事前学習モデルについても概要を紹介します。また、各手法を読者が実際に試せるよう、Google Colaboratory^[7]でのプログラミング例も、同時に示しています。プログラミング例は、理解しやすいよう、また「写経」しやすいよう、長くて 25 行程度に収めることを心がけました。python 等の基本知識が無くとも、とりあえずは動かしてみることが出来るでしょう。

本書の構成は、第 1 章が言語横断的な「概説編」となっており、第 2 章以降で各言語ごとの手法を解説しています。どの章から読み始めても大丈夫ですが、疑問が起これば、必ず第 1 章に戻ってみることをオススメします。なお、それぞれの解析手法に関しては、できるだけ原論文や URL を示すようにしました。また、本書で書き切れなかった諸言語についても、deplacy^[8]のサポートページ^[9]に、Google Colaboratory での使用例を示しました(表 1)。応用を検討する際の一助となれば幸いです。

1.1 Universal Dependencies と CoNLL-U

UD は、書写言語における品詞・形態素属性・依存構造(係り受け関係)を、言語に関わらず記述する手法です。句構造を考慮せずに係り受け関係を記述することで、言語横断性を高めていて、全ての文法構造を単語間のリンクで記述するのが特徴です。

依存構造解析それ自体は、Tesnière の構造的統語論^[10]に源を発し、Мельчук の有向グラフ記述^[11]によって、一応の完成を見た手法です。その最大の特長は、いわゆる動詞中心

^[1]Marie-Catherine de Marneffe, Christopher D. Manning, Joakim Nivre, Daniel Zeman: Universal Dependencies, Computational Linguistics, Vol.47, No.2 (June 2021), pp.255-308.

^[2]<https://universaldependencies.org>において、カレル大学の LINDAT/CLARIN を中心とした 100 以上のグループが、世界中の言語に対して開発中。

^[3]Jacob Devlin, Ming-Wei Chang, Kenton Lee, Kristina Toutanova: BERT: Pre-training of Deep Bidirectional Transformers for Language, Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics (June 2019), Human Language Technologies, Vol.1, pp.4171-4186.

^[4]<https://arxiv.org/abs/1907.11692>

^[5]Pencheng He, Xiaodong Liu, Jianfeng Gao, Weizhu Chen: DeBERTa: Decoding-enhanced Bert With Disentangled Attention, 9th International Conference on Learning Representations (May 2021), Poster 03-2562.

^[6]Benjamin Warner, Antoine Chaffin, Benjamin Clavié, Orion Weller, Oskar Hallström, Said Taghadouini, Alexis Gallagher, Raja Biswas, Faisal Ladhak, Tom Aarsen, Griffin Thomas Adams, Jeremy Howard, Iacopo Poli: Smarter, Better, Faster, Longer: A Modern Bidirectional Encoder for Fast, Memory Efficient, and Long Context Finetuning and Inference, Proceedings of the 63rd Annual Meeting of the Association for Computational Linguistics (July 2025), Vol.1: Long Papers, pp.2526-2547.

^[7]<https://colab.research.google.com>

^[8]安岡孝一: Universal Dependencies にもとづく多言語係り受け可視化ツール deplacy, 人文科学とコンピュータシンポジウム「じんもんこん 2020」論文集 (2020 年 12 月), pp.95-100.

^[9]<https://koichiyasuoka.github.io/deplacy/#templates-for-google-colaboratory>

^[10]Lucien Tesnière: Éléments de Syntaxe Structurale, Paris: C. Klincksieck (1959).

^[11]Igor A. Mel'čuk: Dependency Syntax: Theory and Practice, New York: State University of New York Press (1988).

表 1: deplacy に接続可能な各言語の係り受け解析ツール (2026 年 9 月 15 日現在)

アフリカーンス語	Camphr-Udify, UDPipe 2, Trankit, spaCy-jPTDP, Stanza, spacy-udpipe
アイヌ語	esupar
アラビア語	Trankit, UDPipe 2, spacy-udpipe, Stanza, spaCy-jPTDP
ベラルーシ語	Camphr-Udify, Stanza, Trankit, UDPipe 2
ブルガリア語	Trankit, Camphr-Udify, UDPipe 2, CLASSLA, Stanza, spaCy-jPTDP, spacy-udpipe
カタルーニャ語	Camphr-Udify, Stanza, spacy-udpipe, Trankit, UDPipe 2, spaCy, spaCy-jPTDP
コプト語	esupar, spaCy-Coptic, UDPipe 2, Stanza
チェコ語	Camphr-Udify, spaCy-jPTDP, Trankit, UDPipe 2, spacy-udpipe, Stanza
ウェールズ語	UDPipe 2, Stanza, Camphr-Udify
デンマーク語	Camphr-Udify, UDPipe 2, Stanza, Trankit, spaCy, spacy-udpipe, spaCy-jPTDP
ドイツ語	Camphr-Udify, Stanza, UDPipe 2, spaCy-jPTDP, NLP-Cube, spacy-udpipe, Trankit, spaCy など
ギリシア語	Stanza, UDPipe 2, Trankit, spacy-udpipe, NLP-Cube, spaCy-jPTDP, Camphr-Udify, spaCy など
英語	Camphr-Udify, Stanza, UDPipe 2, NLP-Cube, Trankit, spacy-udpipe, spaCy-jPTDP
スペイン語	Stanza, UDPipe 2, NLP-Cube, spaCy, spacy-udpipe, Trankit, Camphr-Udify, spaCy-jPTDP など
エストニア語	Stanza, spacy-udpipe, EstMalt, spaCy-jPTDP, UDPipe 2, Trankit, Camphr-Udify
バスク語	Trankit, Stanza, spacy-udpipe, Camphr-Udify, spaCy-jPTDP, UDPipe 2
ペルシア語	Stanza, spaCy-jPTDP, spacy-udpipe, Trankit, Camphr-Udify, UDPipe 2 など
スオミ語	Stanza, UDPipe 2, spacy-udpipe, spacy-fi, Trankit, Camphr-Udify, spaCy-jPTDP
フェロー語	UDPipe 2, Stanza
フランス語	spaCy, Stanza, UDPipe 2, Trankit, spacy-udpipe, Camphr-Udify, NLP-Cube, spaCy-jPTDP など
ゲール語 (アイルランド)	Stanza, UDPipe 2, Trankit, spacy-udpipe, spaCy-jPTDP, Camphr-Udify
ゲール語 (スコットランド)	Stanza, UDPipe 2, Trankit
ガリシア語	Camphr-Udify, Stanza, spaCy-jPTDP, Trankit, UDPipe 2
古典ギリシア語	Camphr-Udify, UDPipe 2, spaCy-jPTDP, Stanza, Trankit, spacy-udpipe
ヘブライ語	Trankit, Stanza, spaCy-jPTDP, UDPipe 2, spacy-udpipe
ヒンディー語	UDPipe 2, Trankit, Stanza, spaCy-jPTDP, spacy-udpipe, Camphr-Udify
クロアチア語	UDPipe 2, Trankit, Camphr-Udify, Stanza, CLASSLA, spaCy, spaCy-jPTDP, spacy-udpipe
ハンガリー語	Trankit, UDPipe 2, Camphr-Udify, HuSpacy, spacy-udpipe, Stanza, spaCy-jPTDP, NLP-Cube
アルメニア語	Stanza, Camphr-Udify, Trankit, UDPipe 2, spacy-udpipe
インドネシア語	Stanza, Trankit, spacy-udpipe, Camphr-Udify, UDPipe 2, spaCy-jPTDP
アイスランド語	Stanza, UDPipe 2, is.ud.is.pud
イタリア語	NLP-Cube, Trankit, Camphr-Udify, Stanza, spaCy, UDPipe 2, spaCy-jPTDP, spacy-udpipe
日本語	spaCy-SynCha, spaCy-ChaPAS, UniDic-COMBO, spaCy, UDPipe 2, Stanza, UniDic2UD など
カザフ語	Camphr-Udify, Trankit, Stanza, NLP-Cube, kazdet
韓国語	Camphr-Udify, Stanza, UDPipe 2, Trankit, spaCy-jPTDP, spacy-udpipe, Stanza
ラテン語	Trankit, Stanza, UDPipe 2, spacy-udpipe, Camphr-Udify, spaCy-jPTDP
リトアニア語	spaCy, Trankit, Stanza, Camphr-Udify, UDPipe 2, spacy-udpipe
ラトビア語	Camphr-Udify, Trankit, UDPipe 2, Stanza, spacy-udpipe, spaCy-jPTDP
古典中国語	SuPar-Kanbun, UD-Kanbun, UD-Chinese, Trankit, esupar, GuwenCOMBO, Stanza, UDPipe 2 など
マケドニア語	bert-base-slavic-cyrillic-upos, Camphr-Udify, spaCy
マルタ語	Stanza, UDPipe 2, Camphr-Udify
ノルウェー語 (ブークモール)	Stanza, Camphr-Udify, UDPipe 2, Trankit, spaCy-jPTDP, spaCy
ノルウェー語 (ニーノルスク)	Stanza, spacy-udpipe, spaCy-jPTDP, UDPipe 2, Trankit
オランダ語	spaCy-Alpino, Camphr-Udify, Stanza, UDPipe 2, Trankit, spacy-udpipe, spaCy, spaCy-jPTDP など
ポーランド語	spaCy, Camphr-Udify, Stanza, UDPipe 2, spacy-udpipe, Trankit, spaCy-jPTDP
ポルトガル語	spaCy, Camphr-Udify, Stanza, spaCy-jPTDP, UDPipe 2, spacy-udpipe, Trankit
ルーマニア語	UDPipe 2, Trankit, NLP-Cube, Camphr-Udify, spaCy, Stanza, spaCy-jPTDP, spacy-udpipe
ロシア語	Stanza, Trankit, Camphr-Udify, UDPipe 2, NLP-Cube, spaCy, spacy-udpipe, spaCy-jPTDP など
スロバキア語	Camphr-Udify, UDPipe 2, Trankit, spacy-udpipe, NLP-Cube, Stanza, spaCy-jPTDP
スロベニア語	CLASSLA, UDPipe 2, spacy-udpipe, Camphr-Udify, Trankit, spaCy-jPTDP, Stanza
セルビア語 (キリル)	esupar, Camphr-Udify
セルビア語 (ラテン)	CLASSLA, UDPipe 2, Trankit, esupar, Camphr-Udify, spacy-udpipe, Stanza, spaCy-jPTDP
スウェーデン語	Camphr-Udify, Trankit, Stanza, UDPipe 2, spacy-udpipe, spaCy-jPTDP
タミル語	Camphr-Udify, Trankit, UDPipe 2, spacy-udpipe, Stanza
タイ語	esupar, UDPipe 2, Stanza, spaCy-Thai
タガログ語	Camphr-Udify, ud-tagalog-spacy
トルコ語	spaCy-jPTDP, Stanza, Camphr-Udify, UDPipe 2, spacy-udpipe, Trankit
ウクライナ語	Stanza, NLP-Cube, Camphr-Udify, UDPipe 2, Trankit, spacy-udpipe, spaCy-jPTDP など
ベトナム語	Trankit, Stanza, esupar, spaCy-jPTDP, UDPipe 2, spacy-udpipe など
ウォロフ語	UDPipe 2, Stanza
中国語 (簡化字)	Trankit, Stanza, UDPipe 2, NLP-Cube, esupar, spacy-udpipe, UD-Chinese, spaCy
中国語 (繁体字)	Trankit, Stanza, UDPipe 2, esupar, NLP-Cube, spacy-udpipe, UD-Chinese, spaCy

表 2: CoNLL-U の各フィールド

1. ID: 単語ごとに付与されたインデックスで、文ごとに 1 から始まる整数。縮約語に対しては、単語の範囲を示すのも可。
2. FORM: 語、または、句読記号。
3. LEMMA: 基底形、語幹。
4. UPOS: UD で規定された言語普遍的な品詞タグ (表 3)。
5. XPOS: 言語固有の品詞タグ。
6. FEATS: UD で規定された言語普遍的な形態素属性のリスト。言語固有の拡張も可。
7. HEAD: 当該の単語の係り受け元 ID。係り受け元が無い場合は 0 とする。
8. DEPREL: UD で規定された言語普遍的な係り受けタグ (表 4)。HEAD が 0 の場合は root とする。言語固有の拡張も可。
9. DEPS: 複数の係り受け元を持つ場合、全ての HEAD:DEPREL ペア。
10. MISC: その他のアノテーション。

表 3: UD 品詞タグ (UPOS)

Open class words	Closed class words	Other
ADJ 形容詞 ADV 副詞 INTJ 感嘆詞 NOUN 名詞 PROPN 固有名詞 VERB 動詞	ADP 側置詞 AUX 助動詞 CCONJ 並列接続詞 DET 限定詞 NUM 数詞 PART 接辞 PRON 代名詞 SCONJ 従属接続詞	PUNCT 句読点 SYM 記号 X その他

表 4: UD 係り受けタグ (DEPREL)

	Nominals	Clauses	Modifier Words	Function Words
Core arguments	nsubj 主語 obj 目的語 iobj 間接目的語	csubj 節主語 ccomp 節目的語 xcomp 節補語		
Non-core dependents	obl 斜格補語 vocative 呼称語 expl 形式語 dislocated 外置語	advcl 連用修飾節	advmod 連用修飾語 discourse 談話要素	aux 動詞補助成分 cop 繫辞 mark 標識
Nominal dependents	nmod 体言による連体修飾語 appos 同格 nummod 数量による修飾語	acl 連体修飾節	amod 用言による連体修飾語	det 決定語 clf 類別語 case 格表示
Coordination	MWE	Loose	Special	Other
conj 接続 cc 接続語	fixed 固着 flat 並列 compound 複合	list 細目 parataxis 隣接表現	orphan 親なし goeswith 泣き別れ reparandum 言い損じ	punct 句読点 root 親 dep 未定義

# text = 人莫知其子之惡									
1	人	人	NOUN	n,名詞,人,人	-	3	nsubj	-	SpaceAfter=No
2	莫	莫	ADV	v,副詞,否定,禁止	-	3	advmod	-	SpaceAfter=No
3	知	知	VERB	v,動詞,行為,動作	-	0	root	-	SpaceAfter=No
4	其	其	PRON	n,代名詞,人称,起格	-	5	det	-	SpaceAfter=No
5	子	子	NOUN	n,名詞,人,關係	-	7	nmod	-	SpaceAfter=No
6	之	之	SCONJ	p,助詞,接統,属格	-	5	case	-	SpaceAfter=No
7	惡	惡	NOUN	n,名詞,描写,態度	-	3	obj	-	SpaceAfter=No

# text = 人は其の子の悪しきを知らない									
1	人	人	NOUN	名詞-普通名詞-一般	-	8	nsubj	-	SpaceAfter=No
2	は	は	ADP	助詞-係助詞	-	1	case	-	SpaceAfter=No
3	其の	其の	DET	連体詞	-	4	det	-	SpaceAfter=No
4	子	子	NOUN	名詞-普通名詞-一般	-	6	nsubj	-	SpaceAfter=No
5	の	の	ADP	助詞-格助詞	-	4	case	-	SpaceAfter=No
6	悪しき	悪い	ADJ	形容詞-一般	-	8	ccomp	-	SpaceAfter=No
7	を	を	ADP	助詞-格助詞	-	6	case	-	SpaceAfter=No
8	知ら	知る	VERB	動詞-一般	-	0	root	-	SpaceAfter=No
9	ない	ない	AUX	助動詞	-	8	aux	-	SpaceAfter=No

# text = A man doesn't know the evil of his son									
1	A	a	DET	DT	-	2	det	-	-
2	man	man	NOUN	NN	-	5	nsubj	-	-
3-4	doesn't	-	-	-	-	-	-	-	-
3	does	do	AUX	VBZ	-	5	aux	-	SpaceAfter=No
4	n't	not	PART	RB	-	5	advmod	-	-
5	know	know	VERB	VB	-	0	root	-	-
6	the	the	DET	DT	-	7	det	-	-
7	evil	evil	NOUN	NN	-	5	obj	-	-
8	of	of	ADP	IN	-	10	case	-	-
9	his	he	PRON	PRP\$	-	10	det	-	-
10	son	son	NOUN	NN	-	7	nmod	-	SpaceAfter=No

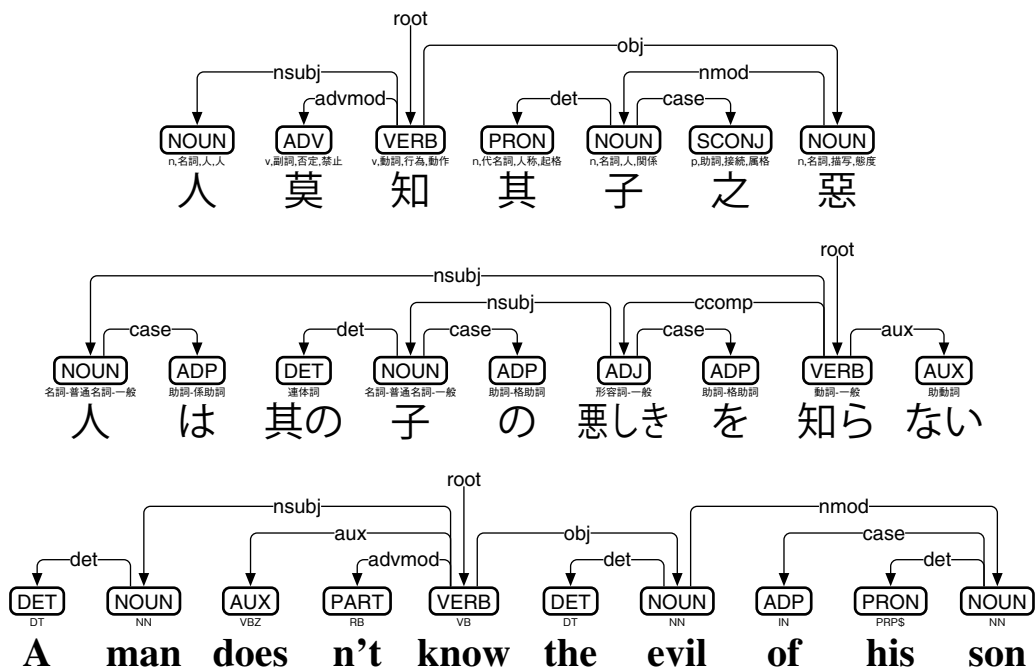


図 1: 漢日英 CoNLL-U と deplacy による可視化

主義によって言語横断的な記述が可能だという点にあり、Мельчук 依存文法をコンピュータ向けに洗練した UD においても、言語に関わらない記述、という特長が前面に押し出されています。UD における文法構造記述は、句構造を考慮せず、全てを単語間のリンクとして表現します。これにより、言語横断的な文法構造記述を可能としているのです。

UD 係り受けコーパスの交換用フォーマットとして、CoNLL-U と呼ばれるタブ区切りテキスト (文字コードは UTF-8^[12]) が規定されています。CoNLL-U の各行は各単語に対応していて、表 2 に示す 10 個のタブ区切りフィールドで構成されます。ID・FORM・LEMMA は、単語そのものに関するフィールドです。UPOS・XPOS・FEATS は、単語の品詞と形態素属性に関するフィールドです。HEAD・DEPREL・DEPS は、単語の係り受けに関するフィールドです。

UD における係り受け関係は、単語間の有向グラフを HEAD と DEPREL で記述します。HEAD は、その単語に入る有向枝のリンク元 ID を示していて、DEPREL は、その有向枝における係り受けタグです。ただし、HEAD が 0 の場合、その枝に入るリンク元は存在しません。リンクの本数は単語の個数に等しく、各リンクのリンク先は、全て互いに異なっています。すなわち、各単語から出るリンクは複数の可能性があります。各単語に入るリンクは 1 つだけなのです。なお、リンクはループしません。

UD の係り受けリンクは、Мельчук 依存文法の後裔にあたり、いわゆる動詞中心主義です。動詞をリンク元として、主語や目的語へとリンクします。修飾関係においては、被修飾語から修飾語へとリンクします。ただし、側置詞 (前置詞や後置詞) を体言の修飾語だとみなす点^[13]が、Мельчук とは異なっています。また、コピュラ文においては動詞中心主義を採らず、補語をリンク元として、主語や繫辞へとリンクします。

UD の言語横断性を示す実際例として、古典中国語 (漢文) 「人莫知其子之惡」と日本語文「人は其の子の悪しきを知らない」と英文「A man doesn't know the evil of his son」の CoNLL-U を、図 1 で比較してみましょう。これらの CoNLL-U は、各言語 UD を横目に手作業で書いたもので、FEATS・DEPS は使用していません。可視化には deplacy を用いています。漢日英で語順は異なっているものの、nsubj (主語) と nmod (体言による連体修飾語)、obj (目的語) と ccomp (節目的語) のリンクが、それぞれ対応しています。

また、コピュラ文に対する UD の例として、露文「Это ручка」と仏文「C'est un stylo」と日本語文「これはペンです」の CoNLL-U を、図 2 で比較してみましょう。露仏日のいずれにおいても、補語から主語へ nsubj が繋がれていて、それぞれに対応していると言えます。なお、仏日においては、繫辞が cop で繋がれています。

1.2 依存文法理論と Universal Dependencies

UD の係り受け構造は、Мельчук の有向グラフ記述を理論的支柱としながらも、実践面では Reed-Kellogg の文法構造^[14]を取り入れています。このあたりについて、簡単に紹介しておきましょう。

^[12]ISO/IEC 10646-1:1993/Amd.2:1996 Information Technology — Universal Multiple-Octet Coded Character Set (UCS) — Part 1: Architecture and Basic Multilingual Plane, Amendment 2: UCS Transformation Format 8 (UTF-8), International Organization for Standardization, Genève (October 15, 1996).

^[13]Joakim Nivre: Towards a Universal Grammar for Natural Language Processing, CICLing 2015: 16th International Conference on Intelligent Text Processing and Computational Linguistics (April 2015), pp.3-16.

^[14]Alonzo Reed and Brainerd Kellogg: Higher Lessons in English: A Work on English Grammar and Composition, New York: Clark & Maynard (1877).

# text = Это ручка									
1	Это	это	PRON	-	-	2	nsubj	-	-
2	ручка	ручка	NOUN	-	-	0	root	-	SpaceAfter=No
# text = C'est un stylo									
1	C'	ce	PRON	-	-	4	nsubj	-	SpaceAfter=No
2	est	être	AUX	-	-	4	cop	-	-
3	un	un	DET	-	-	4	det	-	-
4	stylo	stylo	NOUN	-	-	0	root	-	SpaceAfter=No
# text = これはペンです									
1	これ	此れ	PRON	代名詞	-	3	nsubj	-	SpaceAfter=No
2	は	は	ADP	助詞-係助詞	-	1	case	-	SpaceAfter=No
3	ペン	ペン	NOUN	名詞-普通名詞一般	-	0	root	-	SpaceAfter=No
4	です	です	AUX	助動詞	-	3	cop	-	SpaceAfter=No

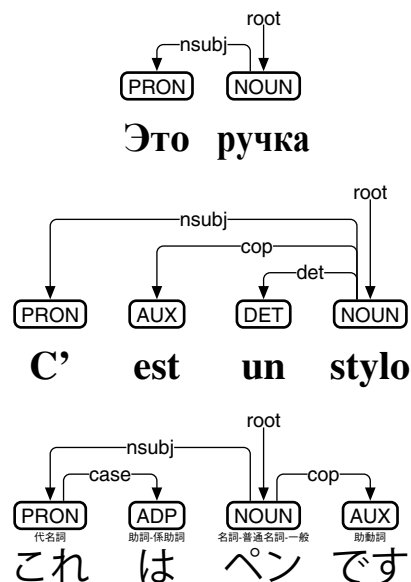


図 2: コピュラ文の CoNLL-U と deplacy による可視化

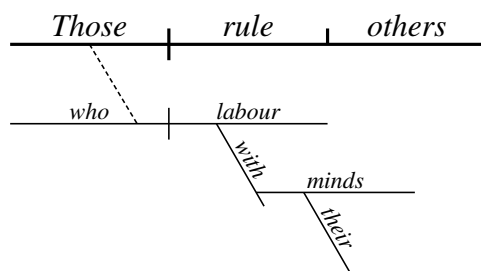


図 3: Reed-Kellogg 文法構造図

Reed-Kellogg 文法構造図は、主語 | 述語 | 目的語 という構造を基本に、英文を視覚化します。主文 (main sentence) を太線で、それ以外の節 (clause) を細線で視覚化した上に、修飾語を斜線で載せて示します。たとえば英文「Those who labour with their minds rule others」に対しては、「Those rule others」という主文の「Those」を、「who labour」という節が修飾していて、その「labour」を「with minds」が修飾し、「minds」を「their」が修飾していることから、図3のようになります。Reed-Kellogg 文法構造図は、英語に特化して設計されているので、必ずしも他の言語には適用できません。

Tesnière は、スラブ諸語の研究を通して、のちに依存文法と呼ばれる言語横断的な記述手法を発案しました。文に現れる語は互いに依存関係にある、というのが、依存文法の基本的な考え方です。語と語の間の依存関係においては、上位項 (régissant) が下位項 (subordonné) を支配し、下位項が上位項に依存します。たとえば英文「Those who labour with their minds rule others」であれば、rule が Those と others を支配し、Those と others が rule に依存します。ただし、転用がおこなわれている場合は、それらの語の間には依存関係は無く、並置する形で表現します。たとえば図4左では、名詞 minds が with によって E 転用 (連用修飾語に転用) されており、あるいは動詞 labour が who によって A 転用 (連体修飾語に転用) されています。

Tesnière の依存文法は、文における語の順序に関わらず、語の依存関係を記述するので、たとえば古典中国語にも応用可能です。すなわち「勞心者治人」であれば、「心」が「勞」に依存し、「勞」が「者」によって O 転用 (体言に転用) され「治」に依存します。これを図4右のように図示すれば、「Those who labour with their minds rule others」との間で、言語をまたいだ比較も可能となるのです。

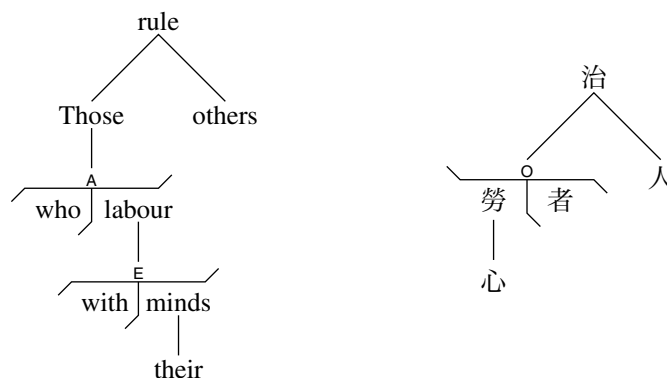


図4: Tesnière 依存文法図

Tesnière の依存文法は、その後、Robinson^[15]や Hudson^[16]らにより、Chomsky 句構造文法^[17]との融合が試みられました。一方Мельчук は、語の依存関係の有向グラフ記述によって依存文法の形式化をおこない、Chomsky 句構造文法と決別しました。

Мельчук の依存文法は、文中の語 X と Y に対し、依存関係 $X \rightarrow Y$ によって文を記述します。 $X \rightarrow Y$ は、X は Y を支配する、あるいは、Y は X に依存する、と読みます。各矢印には、依存関係に関する適切なタグが付与されます。各語が依存する語は高々1つであり、

^[15]Jane J. Robinson: Dependency Structures and Transformational Rules, Language, Vol.46, No.2 (June 1970), pp.259-285.

^[16]Richard Hudson: Word Grammar, New York: Basil Blackwell (1984).

^[17]Noam Chomsky: Aspects of the Theory of Syntax, Cambridge: MIT Press (1965).

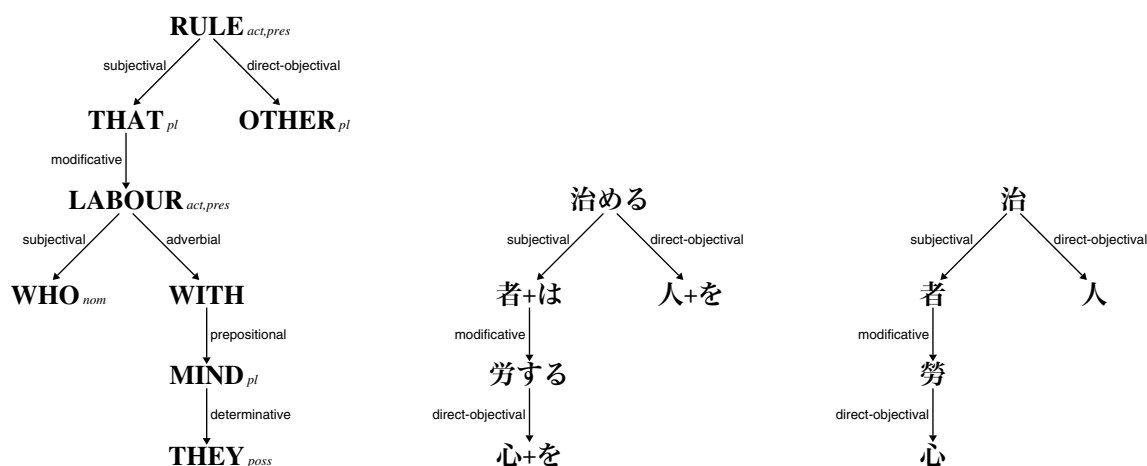


図 5: Мельчук の依存文法による文法構造記述

依存関係はループしません。依存関係は、あくまで語と語の間の記述で、それがたとえ節 (clause) や句 (phrase) に係るものであっても、割り切って語と語の関係として記述します。この「割り切り」の結果として、Мельчук の依存文法は、言語横断的な文法構造記述となっています。英文「Those who labour with their minds rule others」と、日本語文「心を労する者は人を治める」と、古典中国語 (漢文)「勞心者治人」に対して、Мельчук の依存文法による文法構造記述を、図 5 に示します。ただし、タグのうち **subjectival** と **direct-objectival** は、過去には **predicative** と **1st completive** が用いられており、Мельчук 自身も、タグが何種類必要なのかは明らかにしていません^[18]。

UD の係り受け構造は、Мельчук の依存文法におけるタグを、表 4 の 37 種類に限定する野心的な試み、として捉えることができます。図 5 での **subjectival**・**direct-objectival**・**modificative**・**determinative** が、図 6 では **nsubj**・**obj**・**acl**・**det** になった、と考えることもできます。ただ、UD の前身にあたる Stanford Typed Dependencies^[19]が、初期には Bresnan 語彙機能文法^[20]からスタートしており、その後に依存文法へとシフトしたことから、いくつか不思議な概念が紛れ込んでしまっています。その一つが、節 (clause) です。

節という概念を、Мельчук は依存文法から削除したのですが、これを UD は再導入してしまいました。表 4 の **nsubj** と **csbj** は、いずれも主語を表す UD 係り受けタグであり、リンク先が節なら **csbj** を、さもなくば **nsubj** を使います。**obj** と **ccomp** についても同様、**advmod** と **advcl** についても同様、**amod** と **acl** についても同様です。**case** と **mark** については、リンク元が節なら **mark** を、さもなくば **case** を使います^[21]。

英語における節という概念は、一定のコンセンサスが得られていると考えられます。しかし、他の言語において、節という概念は、必ずしも自明ではありません。UD は言語横断的ではあるものの、かなり英語寄りである、という点は意識しておくべきでしょう。

^[18]Alain Polguère, Igor A. Mel'čuk: *Dependency in Linguistic Description*, Amsterdam: John Benjamins (2009).

^[19]Marie-Catherine de Marneffe and Christopher D. Manning: *The Stanford Typed Dependencies Representation*, Coling 2008: Proceedings of the Workshop on Cross-Framework and Cross-Domain Parser Evaluation (August 2008), pp.1-8.

^[20]Joan Bresnan: *Lexical-Functional Syntax*, Malden: Blackwell (2001).

^[21]Marie-Catherine de Marneffe, Timothy Dozat, Natalia Silveira, Katri Haverinen, Filip Ginter, Joakim Nivre, Christopher D. Manning: *Universal Stanford Dependencies: A Cross-Linguistic Typology*, Proceedings of the 9th International Conference on Language Resources and Evaluation (May 2014), pp.4585-4592.

# text = Those who labour with their minds rule others									
1	Those	that	PRON	DT	-	7	nsubj	-	-
2	who	who	PRON	WP	-	3	nsubj	-	-
3	labour	labour	VERB	VBP	-	1	acl	-	-
4	with	with	ADP	IN	-	6	case	-	-
5	their	they	PRON	PRP\$	-	6	det	-	-
6	minds	mind	NOUN	NNS\$	-	3	obl	-	-
7	rule	rule	VERB	VBP	-	0	root	-	-
8	others	other	NOUN	NNS	-	7	obj	-	SpaceAfter=No

# text = 心を労する者は人を治める									
1	心	心	NOUN	名詞-普通名詞-サ変可能	-	3	obj	-	SpaceAfter=No
2	を	を	ADP	助詞-格助詞	-	1	case	-	SpaceAfter=No
3	労する	労する	VERB	動詞-一般	-	4	acl	-	SpaceAfter=No
4	者	者	NOUN	名詞-普通名詞-一般	-	8	nsubj	-	SpaceAfter=No
5	は	は	ADP	助詞-係助詞	-	4	case	-	SpaceAfter=No
6	人	人	NOUN	名詞-普通名詞-一般	-	8	obj	-	SpaceAfter=No
7	を	を	ADP	助詞-格助詞	-	6	case	-	SpaceAfter=No
8	治める	治める	VERB	動詞-一般	-	0	root	-	SpaceAfter=No

# text = 勞心者治人									
1	勞	勞	VERB	v,動詞,描写,境遇	-	3	acl	-	SpaceAfter=No
2	心	心	NOUN	n,名詞,不可讓,身体	-	1	obj	-	SpaceAfter=No
3	者	者	PART	p,助詞,提示,*	-	4	nsubj	-	SpaceAfter=No
4	治	治	VERB	v,動詞,行為,動作	-	0	root	-	SpaceAfter=No
5	人	人	NOUN	n,名詞,人,人	-	4	obj	-	SpaceAfter=No

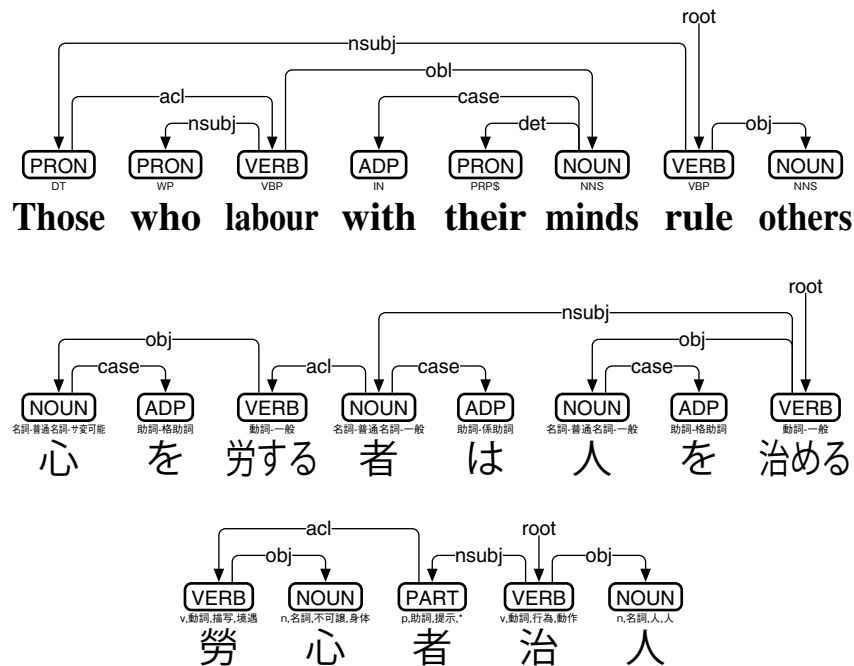


図 6: UD と CoNLL-U による係り受け構造記述

1.3 BERT/RoBERTa/DeBERTa モデルと穴埋めゲーム

Google AI Language が発表した BERT^[3]は、大量の文章を機械に丸暗記させる際にどのような手法が有効なのかを、端的に示した事前学習モデルです。単純に文章を丸暗記するのではなく、文章中に現れる文と文の繋がり、あるいは、文中に現れるトークン(単語もしくは単語より短い文字列)とトークンの繋がりを、効果的に学習できるような工夫が詰め込まれています。

トークンとトークンの繋がりを学習するために用いられるのが、穴埋めゲームです。「This [MASK] a pen.」の [MASK] に対し、「This is a pen.」という答を得るような穴埋めゲーム(図7)を考えてみましょう。bert-base-cased^[22]においては、各トークンの入出力は文字コード(UTF-8)でおこなわれますが、内部的には768次元のベクトル(数値の列)で表現されます。入力側においては、各入力トークンに対応するベクトルをそのまま用いており、出力側においては、各出力トークンに近い^[23]ベクトルになるように学習をおこないます。「This [MASK] a pen.」の例で言えば、「This」「a」「pen」「.」は入力ベクトルをほぼ素通ししつつ、「is」に近い出力ベクトルが得られるよう、内部の記憶素子をいじるわけです。このような穴埋めゲームを大量の文に対しておこなうことで、トークンとトークンの繋がりを学習^[24]できるというのが、BERTの工夫の一つ^[25]です。

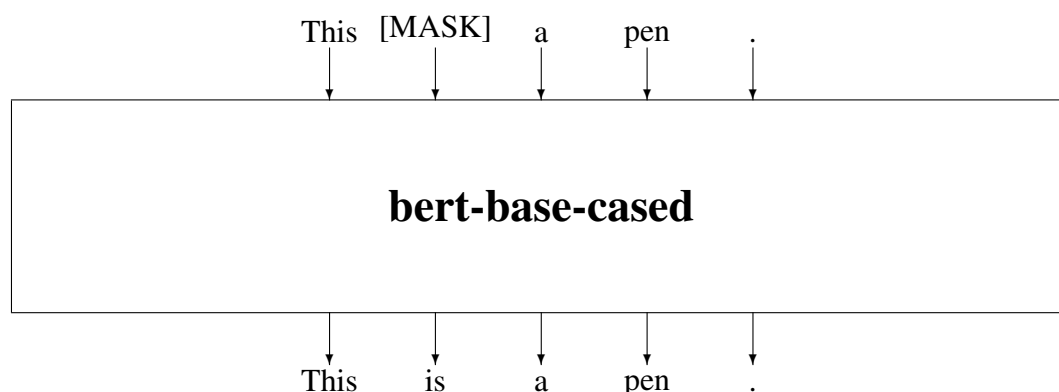


図7: bert-base-cased における穴埋めゲーム

Google Colaboratory 上の Transformers^[26]を使って、bert-base-cased が英文をどのくらい学習できているのか、試してみましょう(図8)。「This [MASK] a pen.」を穴埋めさせると、「was」が48.2%、「is」が40.2%となっており、出力されているベクトルが「was」と「is」の中間で、やや「was」に近いことがわかります。

^[22]<https://huggingface.co/google-bert/bert-base-cased>

^[23]ここでいう「近い」は、コサイン類似度(2つのベクトルの内積を、ベクトルのユークリッド・ノルムで割った値)が1に近い、という意味です。

^[24]bert-base-casedの内部では、穴埋め用の[MASK]に加えて、文頭を表す[CLS]、文の切れ目を表す[SEP]、未知語を表す[UNK]、空きトークンを表す[PAD]なども用いられます。

^[25]BERTのもう一つの工夫は、文と文の繋がりを学習する隣接文ゲームですが、本書では扱いません。

^[26]Thomas Wolf, Lysandre Debut, Victor Sanh, Julien Chaumond, Clement Delangue, Anthony Moi, Pierric Cistac, Tim Rault, Rémi Louf, Morgan Funtowicz, Joe Davison, Sam Shleifer, Patrick von Platen, Clara Ma, Yacine Jernite, Julien Plu, Canwen Xu, Teven Le Scao, Sylvain Gugger, Mariama Drame, Quentin Lhoest, Alexander M. Rush: Transformers: State-of-the-Art Natural Language Processing, Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing (October 2020): System Demonstrations, pp.38-45.

```
!pip install transformers
from transformers import pipeline
fmp=pipeline("fill-mask","google-bert/bert-base-cased")
prd=fmp("This [MASK] a pen.")
print("\n".join("{:8} {:.3f}".format(t["token_str"],t["score"]) for t in prd))
```

was	0.482
is	0.402
from	0.008
had	0.007
has	0.007

図 8: 「This [MASK] a pen.」 に対する穴埋めゲームプログラムと結果

BERT を多言語拡張する際に問題となるのは、モデルの入出力に文という単位を仮定している点です。言語によっては(たとえば古典中国語においては)、文という単位は必ずしも明確ではありません。あるいは SNS での「つぶやき」など、文境界がハッキリしない例は多くあります。これに対し RoBERTa^[4]や DeBERTa^[5]は、文という単位を仮定せず、トークンの並びを用いて、穴埋めゲームのみに特化した学習をおこないます。これにより、文という単位に捉われることなく、大量の文章を丸暗記できるのです。

roberta-classical-chinese-base-char^[27]は、古典中国語における漢字をトークンとみなし、漢字単位での穴埋めゲーム(図 9)を学習したモデルです。各漢字(約 26000 字種)の入出力は UTF-8 ですが、内部的には 768 次元のベクトルで表現します。Google Colaboratory 上の Transformers で「勞 [MASK] 者 治 人」を穴埋めするプログラムと、その結果を図 10 に示します。「心」が 19.9%となっていて、「人」の 32.7%に負けているのがわかります。

BERT・RoBERTa・DeBERTa などの事前学習モデルは、出力部を付け替えることで、様々な用途に応用可能です。出力ベクトルを品詞情報だとみなせば、品詞付与に使えます。係り受けタグだとみなせば、係り受け解析に使えます。UD においては、系列ラベリングや係り受け解析に、事前学習モデルを応用することができます。



図 9: roberta-classical-chinese-base-char における穴埋めゲーム

^[27]<https://huggingface.co/KoichiYasuoka/roberta-classical-chinese-base-char>

```
!pip install transformers
from transformers import pipeline
fmp=pipeline("fill-mask","KoichiYasuoka/roberta-classical-chinese-base-char")
prd=fmp("勞[MASK]者治人")
print("\n".join("{:8} {:.3f}".format(t["token_str"],t["score"]) for t in prd))
```

人	0.327
心	0.199
民	0.129
神	0.054
神	0.047

図 10: 「勞 [MASK] 者治人」 に対する穴埋めゲームプログラムと結果

1.4 GPT 系モデルと言語生成

Open AI が発表した GPT^[28]は、大量の文章を機械に暗記させる際に、言語生成という視点を追及した事前学習モデルです。穴埋めゲームに例えるなら、[MASK] が末尾のみに現れるような穴埋めゲームをおこなっており、次のトークンを予測することに徹しています。トークン予測を連続しておこなうことで、つづきの文章を生成することができます。

gpt2^[29]においては、各トークンの入出力は文字コード (UTF-8) でおこなわれますが、内部的には 768 次元のベクトルで表現されます。入力側においては、各入力トークンに対応するベクトルをそのまま用いており、出力側においては、次のトークンに近い^[23]ベク

```
!pip install transformers
import torch, numpy
from transformers import AutoTokenizer, AutoModelForCausalLM
tkz=AutoTokenizer.from_pretrained("openai-community/gpt2")
mdl=AutoModelForCausalLM.from_pretrained("openai-community/gpt2")
with torch.no_grad():
    m=mdl(**tkz("My name is", return_tensors="pt")).logits.numpy()[0,-1]
    e=numpy.exp(m-numpy.max(m))
    for t in numpy.argsort(-m)[:5]:
        print("{:8} {:.3f}".format(tkz.decode([t]), e[t]/numpy.sum(e)))
```

John	0.011
James	0.008
David	0.008
Michael	0.007
K	0.006

図 11: 「My name is」 の次のトークンを生成するプログラムと結果

^[28]<https://openai.com/index/language-unsupervised>

^[29]<https://huggingface.co/openai-community/gpt2>

トルになるように学習をおこないます。Google Colaboratory 上の Transformers を使って、gpt2 が英文をどう生成するのか試してみましょう (図 11)。「My name is」の次のトークンを生成してみたところ、「John」が 1.1%、「James」と「David」が 0.8%となっていて、男性の名を生成する確率が高めなようです。

GPT 系の言語生成モデルも、出力部を付け替えることで、様々な用途に応用可能です。出力ベクトルを品詞情報だとみなせば、品詞付与に使えます^[30]。係り受けタグだとみなせば、係り受け解析に使えます^[31]。

1.5 Universal Dependencies における係り受け解析

係り受け解析は、UD による言語処理の「キモ」であり、多くの解析アルゴリズムが乱立しています。これらの係り受け解析アルゴリズムを、状態遷移型アルゴリズムと隣接行列型アルゴリズムに分けて、概説します。

arc-swap^[32]に代表される状態遷移型アルゴリズムは、単語列の先頭から末尾に向かって「垣根」(stack-buffer boundary)を移動していく、というイメージで処理をおこないます。「垣根」がおこなう遷移は、以下の 6 種類に定式化されます。

- **Shift** 「垣根」を右に 1 単語分、移動する。
- **Reduce** 「垣根」のすぐ左の単語を除去して、解析結果へ移す。
- **Left-Arc** 「垣根」のすぐ右の単語から、すぐ左の単語へリンクを繋ぐ。
- **Right-Arc** 「垣根」のすぐ左の単語から、すぐ右の単語へリンクを繋ぐ。
- **Unshift** 「垣根」を左に 1 単語分、移動する。
- **Swap** 「垣根」のすぐ右の単語と、すぐ左の単語を入れ替える。

単語が全て **Reduce** されて、「垣根」がポツンと取り残された時点で、状態遷移型アルゴリズムは終了です。解析例を図 12 に示します。解析結果において、入るリンクがない単語に対しては、通常は root を割り当てます。なお、6 種類のうち **Unshift・Swap** を実装しない場合は、リンクに交差がある UD を扱えません。また、現実の実装においては、**Left-Arc** の直後には **Reduce** を、**Right-Arc** の直後には **Shift** と **Reduce** を、それぞれまとめて遷移させる手法が主流です。

Biaffine^[33]に代表される隣接行列型アルゴリズムは、UD 依存構造を有向グラフとみなし、その隣接行列を求めます。たとえば、図 12 右下「Who did you wait for?」の有向グ

ラフであれば
$$\begin{bmatrix} - & - & - & - & \text{case} & - \\ - & - & - & - & - & - \\ - & - & - & - & - & - \\ \text{obl} & \text{aux} & \text{nsbj} & \text{root} & - & \text{punct} \\ - & - & - & - & - & - \\ - & - & - & - & - & - \end{bmatrix}$$
 という 6×6 の隣接行列を求めます (図 13)。

^[30]安岡孝一: GPT 系モデルの系列ラベリングによる品詞付与, 東洋学へのコンピュータ利用, 第 38 回研究セミナー (2024 年 7 月 26 日), pp.3-10.

^[31]安岡孝一: GPT 系言語モデルによる国語研長単位係り受け解析, 人文科学とコンピュータシンポジウム「じんもんこん 2024」論文集 (2024 年 12 月), pp.83-90.

^[32]Joakim Nivre: Non-Projective Dependency Parsing in Expected Linear Time, Proceedings of the Joint Conference of the 47th Annual Meeting of the ACL and the 4th International Joint Conference on Natural Language Processing of the AFNLP (August 2009), pp.351-359.

^[33]Timothy Dozat, Christopher D. Manning: Deep Biaffine Attention for Neural Dependency Parsing, 5th International Conference on Learning Representations (April 2017), C25.

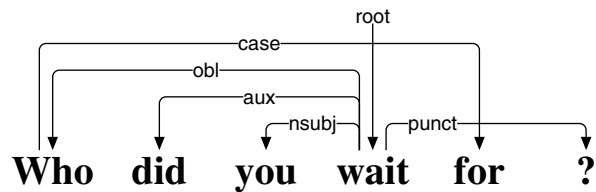
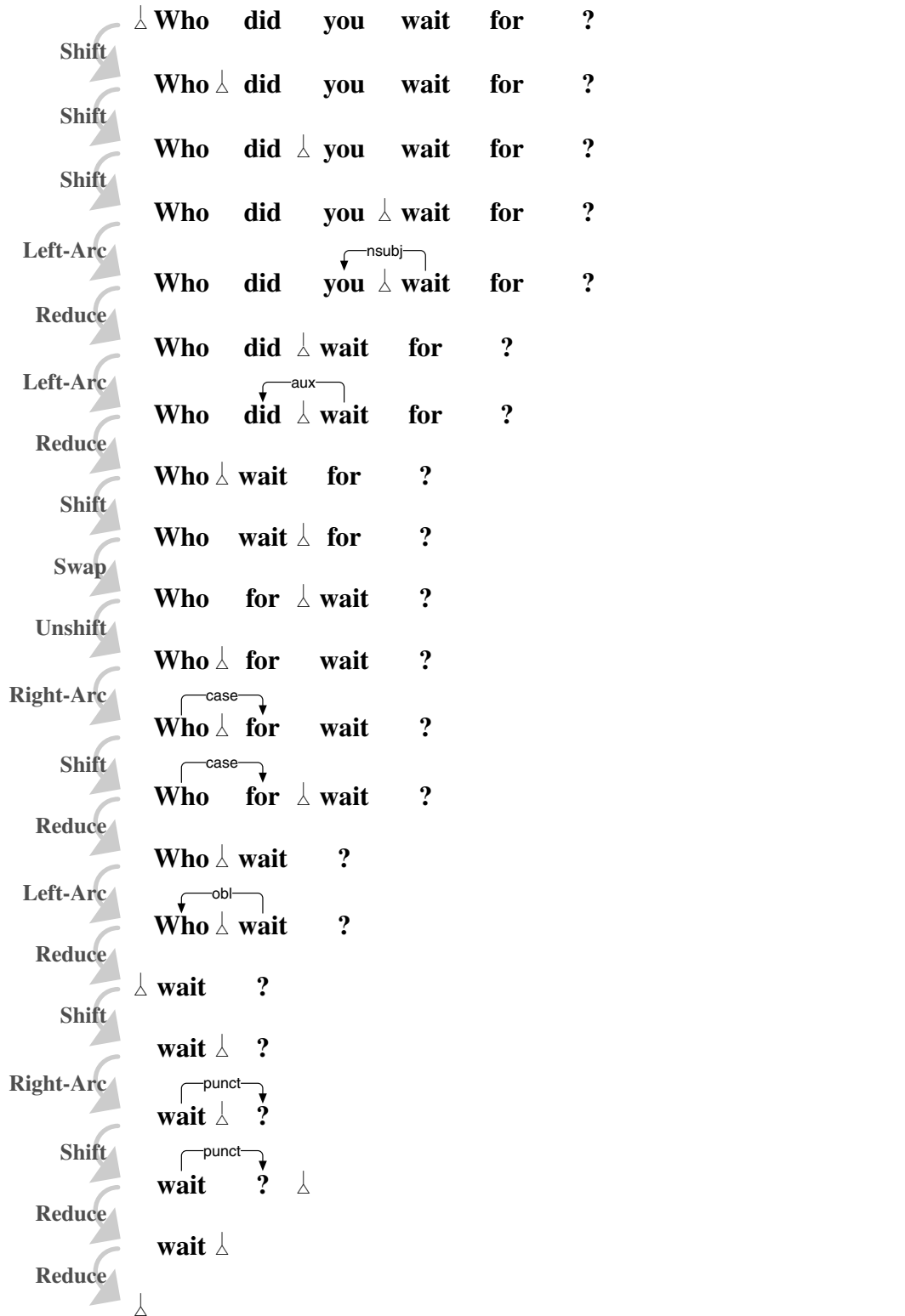


図 12: 状態遷移型アルゴリズムによる「Who did you wait for?」の係り受け解析

具体的には、隣接行列の各列ごとに要素を1つ選びつつ、対角成分上の **root** から1つを選ぶ、というのを確率的におこないます。ただ、ループが発生してしまう場合があるので、適宜 Chu-Liu-Edmonds^{[34][35]}等を使ってループを解消します。

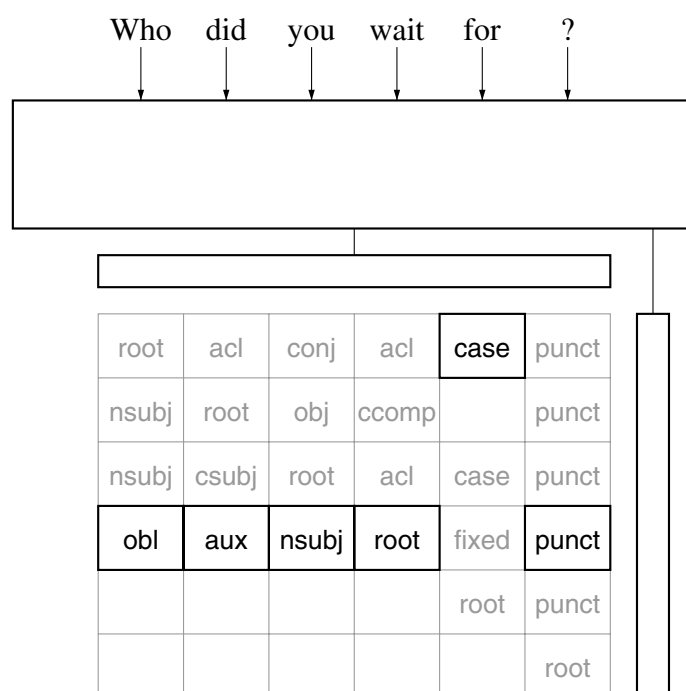


図 13: 隣接行列型アルゴリズムによる「Who did you wait for?」の係り受け解析

状態遷移型アルゴリズムは、BERT・RoBERTa・DeBERTaなどの事前学習モデルを用いることで精度向上が期待できる^[36]のですが、これは実装がかなり難しいのです。一方、隣接行列型アルゴリズムは、事前学習モデルとの相性が良く、実装もそこそこ簡単です。Google Colaboratory 上の Transformers を使って、roberta-base-english-ud-goeswith^[37]で「Who did you wait for?」の隣接行列のロジット^[38]を求めるプログラムと、その結果を図 14 に示します。図 14 の例では、各列ごとにロジット最大の要素を選ぶだけで、「Who did you wait for?」の依存構造グラフを正しく導出できますが、現実の処理では Chu-Liu-Edmonds 等をおこなうことになります。

^[34]Chu Yoeng-jin (朱永津) and Liu Tseng-hong (刘振宏): On the Shortest Arborescence of a Directed Graph, Scientia Sinica, Vol.XIV, No.10 (October 1965), pp.1396-1400.

^[35]Jack Edmonds: Optimum Branchings, Journal of Research of the National Bureau of Standards—B. Mathematics and Mathematical Physics, Vol.71B, No.4 (October-December 1967), pp.233-240.

^[36]Alireza Mohammadshahi, James Henderson: Graph-to-Graph Transformer for Transition-based Dependency Parsing, Findings of the Association for Computational Linguistics: EMNLP 2020 (November 2020), pp.3278-3289.

^[37]<https://huggingface.co/KoichiYasuoka/roberta-base-english-ud-goeswith>

^[38]ラベルの生起確率を p とおくととき、ロジット (対数オッド) は $\log \frac{p}{1-p}$ となります。

```

!pip install transformers
import torch,numpy
from transformers import AutoTokenizer,AutoModelForTokenClassification
brt="KoichiYasuoka/roberta-base-english-ud-goeswith"
txt="Who did you wait for?"
tkz=AutoTokenizer.from_pretrained(brt)
mdl=AutoModelForTokenClassification.from_pretrained(brt)
v,l=tkz(txt,return_offsets_mapping=True),mdl.config.id2label
w=[t for t,(s,e) in zip(v["input_ids"],v["offset_mapping"])] if s<e]
u=[txt[s:e] for s,e in v["offset_mapping"] if s<e]
cls,msk,sep=tkz.cls_token_id,tkz.mask_token_id,tkz.sep_token_id
x=[[cls]+w[:i]+[msk]+w[i+1:]+[sep,j] for i,j in enumerate(w)]
with torch.no_grad():
    m=mdl(input_ids=torch.tensor(x)).logits.numpy()[:-2,: ]
    r=[1 if i==0 else -1 if l[i].endswith("|root") else 0 for i in range(len(l))]
    m+=numpy.where(numpy.add.outer(numpy.identity(m.shape[0]),r)==0,0,-numpy.inf)
    g=mdl.config.label2id["X|_goeswith"]
    r=numpy.tri(m.shape[0])
    for i in range(r.shape[0]):
        for j in range(i+2,r.shape[1]):
            r[i,j]=r[i,j-1] if numpy.argmax(m[i,j-1])==g else 1
    m[:, :,g]+=numpy.where(r==0,0,-numpy.inf)
    d,p=numpy.max(m,axis=2),numpy.argmax(m,axis=2)
    print(" ".join(x[:9].rjust(9) for x in u))
    for i,j in enumerate(u):
        print("\n"+" ".join("{:9.3f}".format(x) for x in d[i])," ",j)
        print(" ".join(l[x].split("|")[-1][:9].rjust(9) for x in p[i]))

```

Who	did	you	wait	for	?	
2.036	2.522	3.163	3.705	12.105	8.971	Who
root	punct	punct	acl:relcl	case	punct	
3.170	0.727	3.536	2.851	8.797	7.458	did
obj	root	punct	xcomp	case	punct	
3.715	1.873	0.808	4.100	9.255	7.630	you
obj	punct	root	parataxis	case	punct	
11.452	14.337	14.130	15.159	8.094	15.684	wait
obl	aux	nsubj	root	compound:	punct	
5.232	1.943	2.694	2.293	2.099	7.492	for
obj	aux	punct	parataxis	root	punct	
1.921	1.544	2.072	1.039	9.437	1.919	?
appos	nmod	punct	nmod	case	root	

図 14: 「Who did you wait for?」に対する隣接行列ロジットの導出

2 古典中国語(漢文)編

2.1 古典中国語UDにおける品詞付与と文切り

古典中国語(漢文)は、単語と単語の間に区切りがなく、文と文の間にも区切りがありません。これが、白文と呼ばれる古典中国語の書写形態であり、傍目には、漢字が連続的に並んでいるだけです。

古典中国語UDにおける品詞付与は、端的に言えば、白文の各漢字にUPOSを付与する作業です。ただし、古典中国語における単語(漢語)は2文字以上のものもあるので、その場合は、複数の漢字をまとめつつUPOSを付与します。Transformers^[26]の系列ラベリングを用いるなら、1文字の漢語に対してはUPOSをラベルとして付与し、2文字以上の漢語に対しては、1文字目に「B-」を付けたUPOSを、2文字目以降に「I-」を付けたUPOSを、それぞれ付与するやり方が考えられます。これにより原理的には、複数の漢字を漢語にまとめつつ、UPOSを付与することができます。

```
!pip install deplacy transformers
class SeqL(object):
    def __init__(self,bert):
        from transformers import pipeline
        self.tagger=pipeline(task="ner",model=bert)
    def __call__(self,text):
        w=[(t["start"],t["end"],t["entity"],t["word"]) for t in self.tagger(text)]
        u,z="# text = "+text.replace("\n"," ")+"\n","_"
        for i,(s,e,p,t) in enumerate(w,1):
            m,q=z if i<len(w) and e<w[i][0] else "SpaceAfter=No",p.split("|")
            f=z if p.find("=")<0 else "|".join(t for t in q if t.find("=")>0)
            u+="\t".join([str(i),t.strip(),z,q[0].replace(".",","),z,f,z,z,z,m])+ "\n"
        return u+"\n"
nlp=SeqL("KoichiYasuoka/roberta-classical-chinese-base-upos")
doc=nlp("子曰學而時習之不亦說乎有朋自遠方來不亦樂乎人不知而不愠不亦君子乎")
import deplacy
deplacy.serve(doc,port=None)
```

NOUN	子	VERB	曰	VERB	學	CCONJ	而	NOUN	時	VERB	習	PRON	之	ADV	不	ADV	亦	VERB	說	PART	乎
VERB	有	NOUN	朋	ADP	自	VERB	遠	NOUN	方	VERB	來	ADV	不	ADV	亦	VERB	樂	PART	乎	NOUN	人
ADV	不	VERB	知	CCONJ	而	ADV	不	VERB	愠	ADV	不	ADV	亦	B-NOUN	君	I-NOUN	子	PART	乎		

図 15: roberta-classical-chinese-base-upos による UPOS 付与

Google Colaboratory 上の Transformers を用いて、roberta-classical-chinese-base-upos^[39] で「子曰學而時習之不亦説乎有朋自遠方來不亦樂乎人不知而不愠不亦君子乎」に UPOS を付与するプログラムと、その結果を図 15 に示します。「君子」が 2 文字の漢語であり、「君」に B-NOUN が、「子」に I-NOUN が付与されています。「君子」以外は漢字 1 字が漢語 1 語と判定されていて、それぞれに UPOS が付与されています。

Transformers の系列ラベリングは、他の用途にも適用できます。roberta-classical-chinese-base-sentence-segmentation^[40] は、古典中国語の文切りに Transformers の系列ラベリングを用いています。具体的には、文頭の漢字に対するラベルを B、文末の漢字に対するラベルを E とし、B・M・E3・E2・E を用いて各文を表します。ただし、1 文字のみで構成される文は、ラベルを S とします。2 文字で構成される文のラベルは、B・E とします。3 文字で構成される文のラベルは、B・E2・E とします。4 文字で構成される文のラベルは、B・E3・E2・E とします。

```
!pip install deplacy transformers
class SeqL(object):
    def __init__(self,bert):
        from transformers import pipeline
        self.tagger=pipeline(task="ner",model=bert)
    def __call__(self,text):
        w=[(t["start"],t["end"],t["entity"],t["word"]) for t in self.tagger(text)]
        u,z="# text = "+text.replace("\n"," ")+"\n","_"
        for i,(s,e,p,t) in enumerate(w,1):
            m,q=z if i<len(w) and e<w[i][0] else "SpaceAfter=No",p.split("|")
            f=z if p.find("=")<0 else "|".join(t for t in q if t.find("=")>0)
            u+="\t".join([str(i),t.strip(),z,q[0].replace(".",","),z,f,z,z,z,m])+"\n"
        return u+"\n"
nlp=SeqL("KoichiYasuoka/roberta-classical-chinese-base-sentence-segmentation")
doc=nlp("子曰學而時習之不亦説乎有朋自遠方來不亦樂乎人不知而不愠不亦君子乎")
import deplacy
deplacy.serve(doc,port=None)
```

ⓑ	ⓔ	ⓑ	Ⓜ	ⓔ3	ⓔ2	ⓔ	ⓑ	ⓔ3	ⓔ2	ⓔ
子	曰	學	而	時	習	之	不	亦	説	乎
ⓑ	Ⓜ	Ⓜ	ⓔ3	ⓔ2	ⓔ	ⓑ	ⓔ3	ⓔ2	ⓔ	ⓑ
有	朋	自	遠	方	來	不	亦	樂	乎	人
Ⓜ	Ⓜ	ⓔ3	ⓔ2	ⓔ	ⓑ	Ⓜ	ⓔ3	ⓔ2	ⓔ	
不	知	而	不	愠	不	亦	君	子	乎	

図 16: roberta-classical-chinese-base-sentence-segmentation による古典中国語の文切り

^[39]<https://huggingface.co/KoichiYasuoka/roberta-classical-chinese-base-upos>

^[40]安岡孝一: Transformers を用いた古典中国語 (漢文) 文切りモデルの製作, 人文科学とコンピュータシンポジウム「じんもんこん 2021」論文集 (2021 年 12 月), pp.104-109.

```

!pip install deplacy transformers
class SeqB(object):
    def __init__(self,bert,trans=[]):
        from transformers import AutoTokenizer,AutoModelForTokenClassification
        import numpy
        self.tokenizer=AutoTokenizer.from_pretrained(bert)
        self.tagger=AutoModelForTokenClassification.from_pretrained(bert)
        x=self.tagger.config.label2id
        self.transition=numpy.full((len(x),len(x)),0 if trans==[] else -numpy.inf)
        for f,t in trans:
            self.transition[x[f],x[t]]=0
    def __call__(self,text):
        import torch,numpy
        v=self.tokenizer(text,return_offsets_mapping=True)
        with torch.no_grad():
            m=self.tagger(torch.tensor([v["input_ids"]])).logits[0].numpy()
        for i in range(m.shape[0]-1,0,-1):
            m[i-1]+=numpy.max(m[i]+self.transition,axis=1)
        j,k=v["offset_mapping"],[numpy.argmax(m[0])]
        for i in range(1,m.shape[0]):
            k.append(numpy.argmax(m[i]+self.transition[k[-1]]))
        w=[(s,e,self.tagger.config.id2label[q]) for (s,e),q in zip(j,k) if s<e]
        u="# text = "+text.replace("\n"," ")+"\n"
        for i,(s,e,p) in enumerate(w,1):
            m="_" if i<len(w) and e<w[i][0] else "SpaceAfter=No"
            u+='\t'.join([str(i),text[s:e],"_",p]+["_"]*5+[m])+"\n"
        return u+"\n"
t=[("S","S"),("S","B"),("B","M"),("B","E3"),("B","E2"),("B","E"),
    ("M","M"),("M","E3"),("E3","E2"),("E2","E"),("E","S"),("E","B")]
nlp=SeqB("KoichiYasuoka/roberta-classical-chinese-base-sentence-segmentation",t)
doc=nlp("子曰學而時習之不亦說乎有朋自遠方來不亦樂乎人不知而不愠不亦君子乎")
import deplacy
deplacy.serve(doc,port=None)

```

ⓑ	ⓔ	ⓑ	Ⓜ	ⓔ3	ⓔ2	ⓔ	ⓑ	ⓔ3	ⓔ2	ⓔ
子	曰	學	而	時	習	之	不	亦	說	乎
ⓑ	Ⓜ	Ⓜ	ⓔ3	ⓔ2	ⓔ	ⓑ	ⓔ3	ⓔ2	ⓔ	ⓑ
有	朋	自	遠	方	來	不	亦	樂	乎	人
Ⓜ	Ⓜ	ⓔ3	ⓔ2	ⓔ	ⓑ	Ⓜ	ⓔ3	ⓔ2	ⓔ	
不	知	而	不	愠	不	亦	君	子	乎	

図 17: 古典中国語の文切りに Bellman-Ford を応用

Google Colaboratory 上の Transformers を用いて、roberta-classical-chinese-base-sentence-segmentation で「子曰學而時習之不亦說乎有朋自遠方來不亦樂乎人不知而不愠不亦君子乎」を文切りするプログラムと、その結果を図 16 に示します。「子曰」「學而時習之」「不亦說乎」「有朋自遠方來」「不亦樂乎」「人不知而不愠」「不亦君子乎」の 7 つの文に切られているのがわかります。ただし、系列ラベリングで文切りをおこなう場合は、Bellman-Ford^[41] や Conditional Random Fields^[42] を応用する方が、より高い精度が得られると考えられます (図 17)。

2.2 古典中国語 UD における係り受け解析

SuPar-Kanbun^[43] は、隣接行列型係り受け解析アルゴリズムを、古典中国語向けに実装しています。UPOS 品詞付与と文切りも同時に実装しており、roberta-classical-chinese-base-char^[27] を元にした「オールインワン設計」です。Google Colaboratory 上の SuPar-Kanbun を用いて、「子曰學而時習之不亦說乎有朋自遠方來不亦樂乎人不知而不愠不亦君子乎」を品詞付与・文切り・係り受け解析するプログラムと、その結果を図 18 に示します。各文の解析結果を、順に見ていくことにしましょう。

「子曰」は、動詞「曰」の主語が「子」という文であり、主語を表す **nsubj** で繋がれています。「學而時習之」は、並列接続詞「而」によって「學」と「習」が並置されていて、「學」と「習」が **conj** で、「習」と「而」が **cc** で繋がれています。「之」は「習」の目的語で、**obj** で繋がれています。「時」は「習」の斜格補語で、本来 **obl** で繋ぐところですが、古典中国語 UD では、**obl** を拡張した **obl:tmod** (時に関する斜格補語) を導入しています。「不亦說乎」は、「不」と「亦」がいずれも「說」を連用修飾していて、いずれも **advmod** で繋がれています。「乎」は文全体を修飾する句末の終助詞 (sentence particle) ですが、これを表すような UD 係り受けタグが表 4 に無いため、**discourse:sp** を導入^[44] しています。「有朋自遠方來」は、「有朋」と「自遠方來」の 2 句から構成されていて、これらが **parataxis** で繋がれています。「朋」は「有」の目的語で、**obj** で繋がれています。「遠方」は「來」の斜格補語で、本来 **obl** で繋ぐところですが、古典中国語 UD では、**obl** を拡張した **obl:lmod** (場所に関する斜格補語) を導入しています。ただし「遠方」の内部では、動詞「遠」(古典中国語 UD には形容詞という分類は無く、動詞として扱います^[45]) が名詞「方」を連体修飾していて、**amod** で繋がれています。なお、前置詞「自」も「方」を修飾しているとみなし、格表示を表す **case** で繋がれています。「不亦樂乎」は「不亦說乎」と同様の依存構造です。「人不知而不愠」は、「而」によって「知」と「愠」が並

^[41]Richard Bellman: On a Routing Problem, Quarterly of Applied Mathematics, Vol.16, No.1 (April 1958), pp.87-90.

^[42]John Lafferty, Andrew McCallum, Fernando Pereira: Conditional Random Fields: Probabilistic Models for Segmenting and Labeling Sequence Data, Proceedings of the 18th International Conference on Machine Learning (June 2001), pp.282-289.

^[43]安岡孝一, ウィッテルン クリスティアン, 守岡知彦, 池田巧, 山崎直樹, 二階堂善弘, 鈴木慎吾, 師茂樹, 藤田一乗: 古典中国語 (漢文) Universal Dependencies とその応用, 情報処理学会論文誌, Vol.63, No.2 (2022 年 2 月), pp.355-363.

^[44]Herman Leung, Rafaël Poirer, Tak-sum Wong, Xinying Chen, Kim Gerdes and John Lee: Developing Universal Dependencies for Mandarin Chinese, 12th Workshop on Asian Language Resources (December 2016), pp.20-29.

^[45]安岡孝一, ウィッテルン クリスティアン, 守岡知彦, 池田巧, 山崎直樹, 二階堂善弘, 鈴木慎吾, 師茂樹: 古典中国語 (漢文) の形態素解析とその応用, 情報処理学会論文誌, Vol.59, No.2 (2018 年 2 月), pp.323-331.

```

!pip install suparkanbun
import suparkanbun
nlp=suparkanbun.load(Danku=True)
doc=nlp("子曰學而時習之不亦說乎有朋自遠方來不亦樂乎人不知而不愠不亦君子乎")
import deplacy
deplacy.serve(doc,port=None)

```

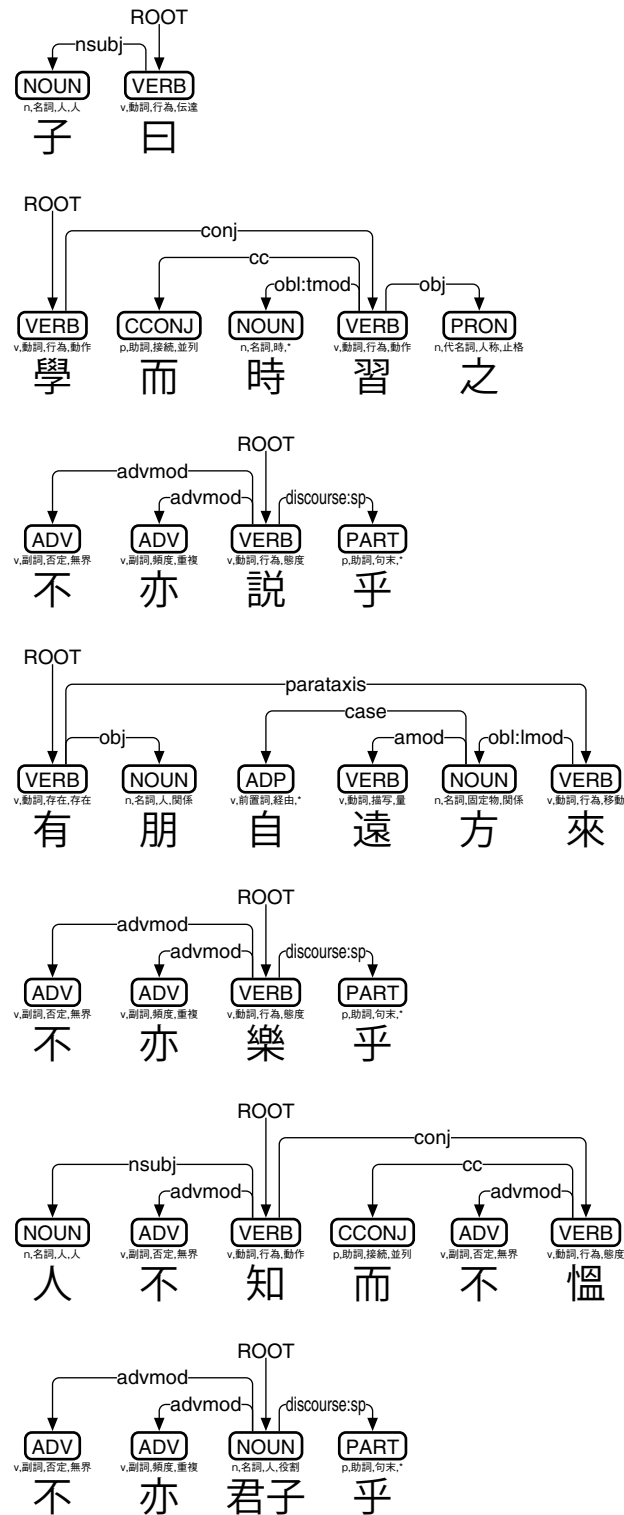


図 18: SuPar-Kanbun による古典中国語の係り受け解析

置されていて、「知」と「慍」が **conj** で、「慍」と「而」が **cc** で繋がられています。「人」は「知」と「慍」の両方の主語だと考えられるのですが、近い方の「知」に **nsubj** で繋がられています。また、「知」と「慍」には、それぞれ否定の「不」が連用修飾していて、いずれも **advmod** で繋がられています。「不亦君子乎」は、「不亦説乎」や「不亦樂乎」と類似した依存構造ですが、「君子」が名詞でありコピュラ文です(ただし主語も繫辞もあります)。

このような形で SuPar-Kanbun は、漢字の列である白文に品詞付与をおこない、さらに漢語の列に係り受け解析をおこなうことで、漢文の依存構造を抽出します。なお、古典中国語 UD の係り受けタグには、表 4 と図 18 に加え、受動文の主語を表す **nsubj:pass**、重畳を表す **compound:redup**、動詞並列を表す **flat:vv**、古典中国語以外を表す **flat:foreign** も用いられます^[46]。

2.3 古典中国語ミニ DeBERTa モデルの製作

Google Colaboratory 上の esupar^[47]で、古典中国語ミニ DeBERTa モデルを製作してみましょう。手順を図 19～23 に示します。

最初に、訓練のための白文を準備しましょう(図 19)。train.txt の白文の材料として、こ

```
!pip install esupar accelerate
import os
url="https://github.com/UniversalDependencies/UD_Classical_Chinese-Kyoto"
d=os.path.basename(url)
!test -d {d} || git clone --depth=1 {url}
!for F in train dev test ; do cp {d}/*-$$F.conllu $$F.conllu ; done
!sed -n "s/^# text = //p" *.conllu > train.txt
```

図 19: 古典中国語ミニモデル製作のための準備

```
from transformers import BertTokenizerFast
from tokenizers.pre_tokenizers import Sequence,Whitespace,Split
from tokenizers import Regex
s=["[CLS]","[PAD]","[SEP]","[UNK]","[MASK]"]
with open("train.txt","r",encoding="utf-8") as r:
    v=set(c for c in r.read() if not c.isspace())
with open("vocab.txt","w",encoding="utf-8") as w:
    print("\n".join(s+sorted(v)),file=w)
tkz=BertTokenizerFast(vocab_file="vocab.txt",never_split=s,
    do_lower_case=False,strip_accents=False,tokenize_chinese_chars=True)
b=tkz.backend_tokenizer
b.pre_tokenizer=Sequence([Whitespace(),Split(Regex("."),"isolated")])
b.decoder.prefix=b.model.continuing_subword_prefix=""
tkz.save_pretrained("my-dir/tokenizer-lzh")
```

図 20: 古典中国語向け単文字トークナイザの作成

^[46]<https://universaldependencies.org/lzh/#syntax>

^[47]<https://github.com/KoichiYasuoka/esupar>

```

from transformers import (AutoTokenizer,DebertaV2Config,DebertaV2ForMaskedLM,
    DataCollatorForLanguageModeling,TrainingArguments,Trainer)
tkz=AutoTokenizer.from_pretrained("my-dir/tokenizer-lzh",model_max_length=128)
tkz.backend_tokenizer.model.max_input_chars_per_word=tkz.model_max_length
t=tkz.convert_tokens_to_ids(["[CLS]","[PAD]","[SEP]","[UNK]","[MASK]"])
cfg=DebertaV2Config(hidden_size=256,num_hidden_layers=12,num_attention_heads=4,
    intermediate_size=768,relative_attention=True,position_biased_input=False,
    pos_att_type=["p2c","c2p"],max_position_embeddings=tkz.model_max_length,
    vocab_size=len(tkz),tokenizer_class=type(tkz).__name__,
    bos_token_id=t[0],pad_token_id=t[1],eos_token_id=t[2])
arg=TrainingArguments(report_to="none",output_dir="/tmp",save_total_limit=2,
    num_train_epochs=3,per_device_train_batch_size=64)
class ReadLineDS(object):
    def __init__(self,file,tokenizer):
        self.tokenizer=tokenizer
        with open(file,"r",encoding="utf-8") as r:
            self.lines=[s.strip() for s in r if s.strip()!=""]
    __len__=lambda self:len(self.lines)
    __getitem__=lambda self,i:self.tokenizer(self.lines[i],truncation=True,
        add_special_tokens=True,max_length=self.tokenizer.model_max_length-2)
trn=Trainer(args=arg,data_collator=DataCollatorForLanguageModeling(tkz),
    model=DebertaV2ForMaskedLM(cfg),train_dataset=ReadLineDS("train.txt",tkz))
trn.train()
trn.save_model("my-dir/deberta-lzh")
tkz.save_pretrained("my-dir/deberta-lzh")

```

図 21: 古典中国語ミニ DeBERTa モデルの作成

```

import torch
from transformers import AutoTokenizer,AutoModelForMaskedLM
from esupar.tradify import tradify
tkz=AutoTokenizer.from_pretrained("my-dir/deberta-lzh")
mdl=AutoModelForMaskedLM.from_pretrained("my-dir/deberta-lzh")
c=[(k,v) for k,v in tradify.items() if tkz.add_tokens([k,v])==1]
e=mdl.resize_token_embeddings(len(tkz))
with torch.no_grad():
    for k,v in c:
        t=sorted(tkz.convert_tokens_to_ids([k,v]))
        e.weight[t[1],:]=e.weight[t[0],:]
mdl.set_input_embeddings(e)
mdl.save_pretrained("my-dir/deberta-lzh-ext")
tkz.save_pretrained("my-dir/deberta-lzh-ext")

```

図 22: 古典中国語ミニ DeBERTa モデルの拡張 (異体字追加)

```

!python -m esupar.train my-dir/deberta-lzh-ext my-dir/deberta-lzh-upos .

```

図 23: 古典中国語向け品詞付与・係り受け解析ミニモデルの製作

ここでは UD_Classical_Chinese-Kyoto^[48]を元にしていますが、好みに応じ、他の白文を増量するのもいいでしょう。なお、UD_Classical_Chinese-Kyoto の各 CoNLL-U ファイルは、品詞付与・係り受け解析の学習にも用います。

続けて、入力文字列を単文字に分解するためのトークナイザを、my-dir/tokenizer-lzh に作成しましょう。図 20 では、BertTokenizerFast という Transformers のトークナイザを使って、単文字トークナイザを作成しています。

次に train.txt の白文から、DeBERTa モデルを my-dir/deberta-lzh に作成しましょう。図 21 では、入出力幅 128 トークン・入出力ベクトル 256 次元・深さ 12 層・アテンションヘッド 4 個・中間ベクトル 768 次元という、やや小さめのモデルにしており、CPU なら 1 時間程度、GPU なら 10 分程度で作成できます。ただし、この DeBERTa モデルは繁体字で作られており、日本の常用漢字や中国の簡化字は含まれていません。そこで、これらの異体字に対しベクトルのコピーをおこない、新たなモデルを my-dir/deberta-lzh-ext に保存しましょう (図 22)。

最後に、my-dir/deberta-lzh-ext と UD_Classical_Chinese-Kyoto から、品詞付与・係り受け解析モデル my-dir/deberta-lzh-upos を製作しましょう (図 23)。ただ、隣接行列型アルゴリズムの学習は収束が遅く、GPU でも 3 時間程度かかってしまいます。

うまく my-dir/deberta-lzh-upos が製作できたかどうか、esupar でテストしてみましょう (図 24)。繁体字のみならず、日本の常用漢字や中国の簡化字で書かれた漢文も、品詞付与・係り受け解析できるはずなので、ぜひ試してみてください。

ちなみに図 25 は、my-dir/deberta-lzh-ext と UD_Classical_Chinese-Kyoto から、ミニ文切りモデル my-dir/deberta-lzh-seg を作成 (ファインチューニング) するプログラムです。ただし、ミニモデルでは文切りの精度が上がらないので、もっと大きなモデルを作成した方がいいでしょう。

```
!pip install esupar
import esupar
nlp=esupar.load("my-dir/deberta-lzh-upos")
doc=nlp("劳心者治人")
import deplacy
deplacy.serve(doc,port=None)
```

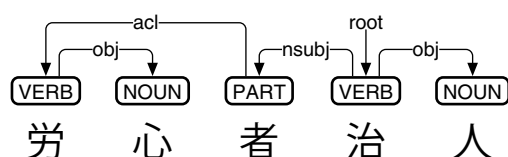


図 24: 古典中国語向け品詞付与・係り受け解析ミニモデルのテスト

^[48]Koichi Yasuoka: Universal Dependencies Treebank of the Four Books in Classical Chinese, DADH2019: 10th International Conference of Digital Archives and Digital Humanities (December 2019), pp.20-28.

```

from transformers import (AutoTokenizer, AutoModelForTokenClassification,
    AutoConfig, DataCollatorForTokenClassification, TrainingArguments, Trainer)
src, lid="my-dir/deberta-lzh-ext", {"B":0, "E":1, "E2":2, "E3":3, "M":4, "S":5}
tkz, cfg=AutoTokenizer.from_pretrained(src), AutoConfig.from_pretrained(src,
    num_labels=len(lid), label2id=lid, id2label={i:1 for l,i in lid.items()})
wd=cfg.max_position_embeddings-3
dat=[{"input_ids": [tkz.cls_token_id], "labels": [5]}]
with open("train.conllu", "r", encoding="utf-8") as r:
    for t in [s[9:].strip() for s in r if s.startswith("# text = ")]:
        i=tkz(t, add_special_tokens=False)["input_ids"]
        j=[0 if len(i)>1 else 5]+[min(k,4) for k in range(len(i)-1,0,-1)]
        if len(dat[-1]["input_ids"])+len(i)>wd:
            dat.append({"input_ids": [tkz.cls_token_id]+i[0:wd], "labels": [5]+j[0:wd]})
        elif len(i)>0:
            dat[-1]["input_ids"].extend(i)
            dat[-1]["labels"].extend(j)
arg=TrainingArguments(report_to="none", output_dir="/tmp", save_total_limit=2,
    num_train_epochs=3, per_device_train_batch_size=64)
dcl=DataCollatorForTokenClassification(tkz)
trn=Trainer(args=arg, data_collator=dcl, train_dataset=dat,
    model=AutoModelForTokenClassification.from_pretrained(src, config=cfg))
trn.train()
trn.save_model("my-dir/deberta-lzh-seg")
tkz.save_pretrained("my-dir/deberta-lzh-seg")

```

図 25: 古典中国語文切りミニモデルの作成

3 日本語編

書写言語としての日本語は、単語と単語の間に区切りがありません。そもそも、単語の長さに決まりがないので、様々な単語長が使われている^[49]のです(図 26)。日本語 UD では、これらの単語長のうち、国語研短単位と国語研長単位がサポートされています。その一方で、日本語の事前学習モデルのトークン長は、国語研短単位(あるいはその細分化)や文字単位のものが多いのですが、図 26 以外のトークン長を採用しているモデルもあり、かなり混沌としています。

国語研短単位	全 学年 に わたっ て 小 学校 の 国語 の 教科 書 に 大量 の 挿し絵 が 用い られ て いる
国語研長単位	全学年 にわたって 小学校 の 国語 の 教科書 に 大量 の 挿し絵 が 用い られ ている
文節	全学年にわたって 小学校の 国語の 教科書に 大量の 挿し絵が 用いられている

図 26: 国語研短単位・国語研長単位・文節

3.1 国語研短単位にもとづく品詞付与と係り受け解析

SuPar-UniDic^[50]は、青空文庫(+Wikipedia) BERT モデル^[51]をベースに、隣接行列型係り受け解析アルゴリズムを、国語研短単位向けに実装しています。形態素解析(単語切り・XPOS 品詞付与)部には、MeCab^[52]と UniDic^[53]を流用しています。Google Colaboratory 上の SuPar-UniDic を用いて、「全学年にわたって小学校の国語の教科書に大量の挿し絵が用いられている」を品詞付与・係り受け解析するプログラムと、その結果を図 27 に示します。解析結果を見ていくことにしましょう。

国語研短単位 UD では、「全学年」「小学校」「教科書」が複数の単語に分かれていて、**compound** で繋がれているのが特徴的です。「小学校の国語の教科書」は、名詞が名詞を修飾していて **nmod** で繋がれていると同時に、「の」が **case** で繋がれています。「大量の挿し絵」も同様です。「全学年」は「わたって」の二格であり、**obl** で繋がれていると同時に、「に」が **case** で繋がれています。「挿し絵」は「用いられて」のガ格であり、**nsubj** で繋がれていると同時に、「が」が **case** で繋がれています。「全学年にわたって」は「挿

^[49]Mai Omura, Aya Wakasa, Masayuki Asahara: Word Delimitation Issues in UD Japanese, Proceedings of the 5th Workshop on Universal Dependencies (December 2021), pp.142-150.

^[50]安岡孝一: 世界の Universal Dependencies と係り受け解析ツール群, 第 3 回 Universal Dependencies 公開研究会 (2021 年 6 月).

^[51]<https://github.com/akirakubo/bert-japanese-aozora>

^[52]Taku Kudo, Kaoru Yamamoto, Yuji Matsumoto: Applying Conditional Random Fields to Japanese Morphological Analysis, Proceedings of the 2004 Conference on Empirical Methods in Natural Language Processing (July 2004), pp.230-237.

^[53]伝康晴, 小木曾智信, 小椋秀樹, 山田篤, 峯松信明, 内元清貴, 小磯花絵: コーパス日本語学のための言語資源: 形態素解析用電子化辞書の開発とその応用, 日本語科学, 第 22 号 (2007 年 10 月), pp.101-123.

```
!pip install suparunidic
import suparunidic
nlp=suparunidic.load("gendai")
doc=nlp("全学年にわたって小学校の国語の教科書に大量の挿し絵が用いられている")
import deplacy
deplacy.serve(doc,port=None)
```

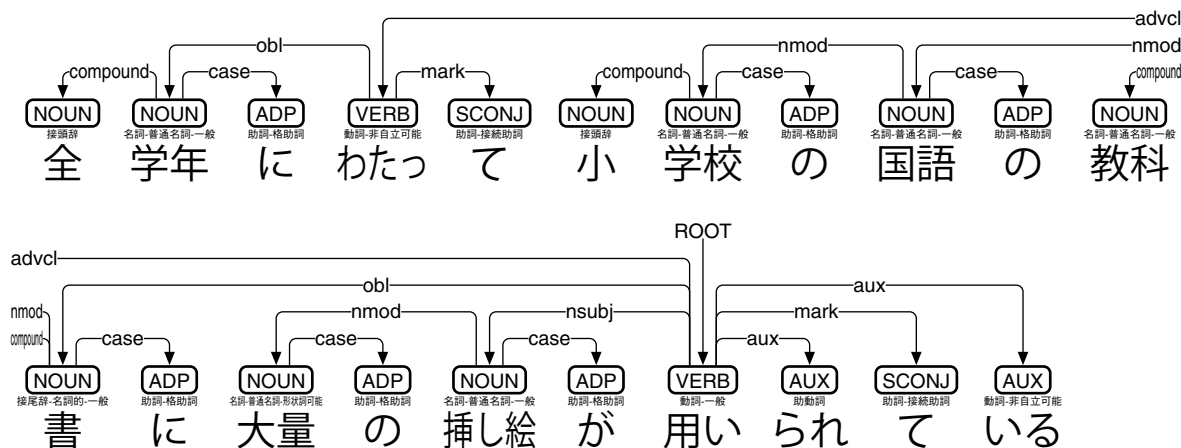


図 27: SuPar-UniDic による日本語係り受け解析 (国語研短単位)

し絵が用いられて」の連用修飾節であり、advcl で繋がられています。「いる」は、UniDic 品詞 (XPOS) が「動詞-非自立可能」なのですが、ここでは非自立だとみなして UPOS を AUX にすると同時に、aux で繋がられています。

国語研短単位 UD では、XPOS と UPOS が微妙に乖離しています。XPOS は UniDic 品詞を流用しているのですが、UPOS の品詞付与ルールとはズレているためです。極端な例として「難儀な難儀は難儀する」(図 28) を見てみましょう。「難儀」の UniDic 品詞は「名詞-普通名詞-サ変形状詞可能」ですが、「難儀な」は形状詞 (形容動詞) として扱うべきであり、UPOS は ADJ が妥当です。「難儀する」はサ変動詞として扱うべきで、UPOS は VERB が妥当です。ただし「難儀な」も「難儀する」も、国語研短単位では 2 つの単語に分かれてしまうので、「難儀」の UPOS に、それぞれ ADJ と VERB を付与しています。

```
!pip install suparunidic
import suparunidic
nlp=suparunidic.load("gendai")
doc=nlp("難儀な難儀は難儀する")
import deplacy
deplacy.serve(doc,port=None)
```

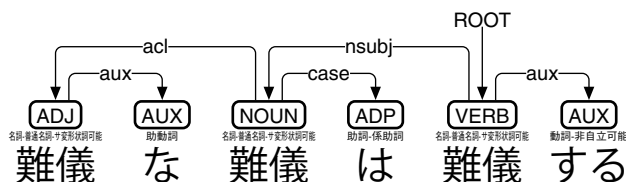


図 28: 国語研短単位 UD における XPOS (UniDic 品詞) と UPOS の関係

3.2 国語研短単位ミニ DeBERTa モデルの製作

Google Colaboratory 上の esupar^[47]で、国語研短単位ミニ DeBERTa モデルを製作してみましょう。手順を図 29～32 に示します。

最初に、訓練のための文章を準備しましょう (図 29)。train.txt の文章の材料として、ここでは UD_Japanese-GSD^[54]と Wikitext-JA^[55]を元にしていますが、好みに応じ、他の日本語文を増量するのもいいでしょう。なお、UD_Japanese-GSD の各 CoNLL-U ファイルは、品詞付与・係り受け解析の学習にも用います。

```
!pip install esupar fugashi unidic-lite accelerate
import os,urllib.request
url="https://github.com/UniversalDependencies/UD_Japanese-GSD"
d=os.path.basename(url)
!test -d {d} || git clone --depth=1 {url}
!for F in train dev test ; do cp {d}/*-$$F.conllu $$F.conllu ; done
!sed -n "s/^# text = //p" *.conllu > train.txt
url="http://www.lsta.media.kyoto-u.ac.jp/resource/data/wikitext-ja/"
with open("train.txt","a",encoding="utf-8") as w:
    for t in ["Featured_Contents.txt","Good_Contents.txt"]:
        with urllib.request.urlopen(url+t) as r:
            print(r.read().decode("utf-8").replace("。","。 \n"),file=w)
```

図 29: 国語研短単位ミニモデル製作のための準備

続けて、入力文を国語研短単位に分解するためのトークナイザを、my-dir/tokenizer-suw に作成しましょう。図 30 では、fugashi^[56]と unidic-lite^[57]で train.txt を国語研短単位に区切り、それを BertJapaneseTokenizer という Transformers^[26]のトークナイザに組み上げて、国語研短単位トークナイザを作成しています。語彙数 (vocab_size) を 8000 に、異なり

```
!fugashi -Owakati < train.txt > token.txt
from transformers import BertJapaneseTokenizer
from tokenizers import BertWordPieceTokenizer
s=["[CLS]","[PAD]","[SEP]","[UNK]","[MASK]"]
wpt=BertWordPieceTokenizer(lowercase=False,strip_accents=False,
    handle_chinese_chars=False)
wpt.train(files=["token.txt"],vocab_size=8000,limit_alphabet=3000,
    special_tokens=s)
wpt.save_model(".")
tkz=BertJapaneseTokenizer(vocab_file="vocab.txt",word_tokenizer_type="mecab",
    mecab_kwargs={"mecab_dic":"unidic-lite"},do_lower_case=False,never_split=s)
tkz.save_pretrained("my-dir/tokenizer-suw")
```

図 30: 国語研短単位トークナイザの作成

^[54]松田寛, 若狭絢, 山下華代, 大村舞, 浅原正幸: UD Japanese GSD の再整備と固有表現情報付与, 言語処理学会第 26 回年次大会発表論文集 (2020 年 3 月), pp.133-136.

^[55]<http://www.lsta.media.kyoto-u.ac.jp/resource/data/wikitext-ja/>

^[56]<https://github.com/polm/fugashi>

^[57]<https://github.com/polm/unidic-lite>

```

from transformers import (AutoTokenizer,DebertaV2Config,DebertaV2ForMaskedLM,
    DataCollatorForLanguageModeling,TrainingArguments,Trainer)
tkz=AutoTokenizer.from_pretrained("my-dir/tokenizer-suw",model_max_length=128)
t=tkz.convert_tokens_to_ids(["[CLS]","[PAD]","[SEP]","[UNK]","[MASK]"])
cfg=DebertaV2Config(hidden_size=256,num_hidden_layers=12,num_attention_heads=4,
    intermediate_size=768,relative_attention=True,position_biased_input=False,
    pos_att_type=["p2c","c2p"],max_position_embeddings=tkz.model_max_length,
    vocab_size=len(tkz),tokenizer_class=type(tkz).__name__,
    bos_token_id=t[0],pad_token_id=t[1],eos_token_id=t[2])
arg=TrainingArguments(report_to="none",output_dir="/tmp",save_total_limit=2,
    num_train_epochs=3,per_device_train_batch_size=64)
class ReadLineDS(object):
    def __init__(self,file,tokenizer):
        self.tokenizer=tokenizer
        with open(file,"r",encoding="utf-8") as r:
            self.lines=[s.strip() for s in r if s.strip()!=""]
    __len__=lambda self:len(self.lines)
    __getitem__=lambda self,i:self.tokenizer(self.lines[i],truncation=True,
        add_special_tokens=True,max_length=self.tokenizer.model_max_length-2)
trn=Trainer(args=arg,data_collator=DataCollatorForLanguageModeling(tkz),
    model=DebertaV2ForMaskedLM(cfg),train_dataset=ReadLineDS("train.txt",tkz))
trn.train()
trn.save_model("my-dir/deberta-suw")
tkz.save_pretrained("my-dir/deberta-suw")

```

図 31: 国語研短単位ミニ DeBERTa モデルの作成

```
!python -m esupar.train my-dir/deberta-suw my-dir/deberta-suw-upos .
```

図 32: 国語研短単位向け品詞付与・係り受け解析ミニモデルの製作

字数(limit_alphabet)を 3000 に絞っていますが、もう少し大きい方がいいかもしれません。

次に train.txt の文章から、DeBERTa モデルを my-dir/deberta-suw に作成しましょう。図 31 では、入出力幅 128 トークン・入出力ベクトル 256 次元・深さ 12 層・アテンションヘッド 4 個・中間ベクトル 768 次元という、やや小さめのモデルにしていますが、GPU でも 3 時間程度かかってしまいます。

最後に、my-dir/deberta-suw と UD_Japanese-GSD から、品詞付与・係り受け解析モデル my-dir/deberta-suw-upos を製作しましょう(図 32)。ただ、隣接行列型アルゴリズムの学習は収束が遅く、GPU でも 2 時間程度かかってしまいます。

うまく my-dir/deberta-suw-upos が製作できたかどうか、esupar でテストしてみましょう。図 33 の例では、「にわたって」「ている」のリンクに **fixed** が現れていて、図 27 とは異なる結果になっているようです。

```
!pip install esupar fugashi unidic-lite
import esupar
nlp=esupar.load("my-dir/deberta-suw-upos")
doc=nlp("全学年にわたって小学校の国語の教科書に大量の挿し絵が用いられている")
import deplacy
deplacy.serve(doc,port=None)
```

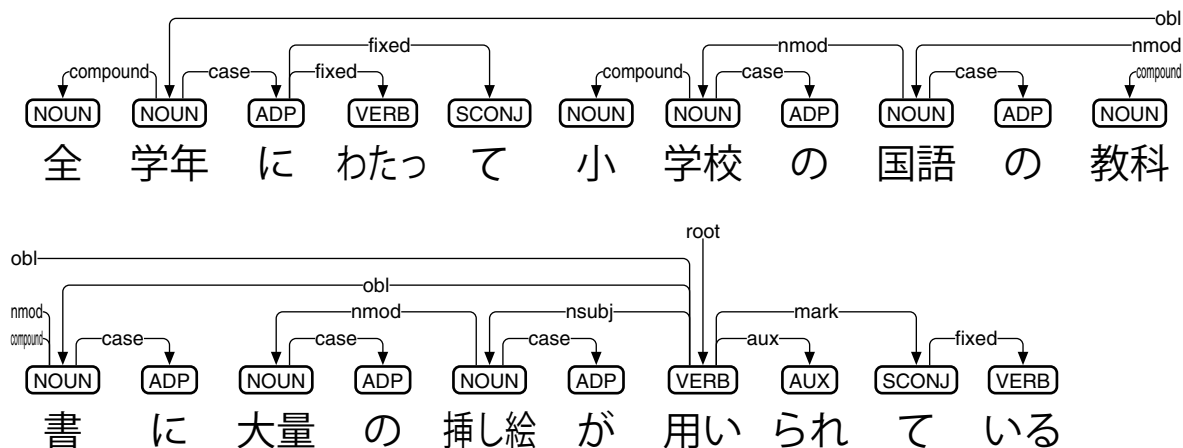


図 33: 国語研短単位向け品詞付与・係り受け解析ミニモデルのテスト

3.3 国語研長単位にもとづく品詞付与と係り受け解析

国語研長単位トークナイザ Japanese-LUW-Tokenizer^[58]と、青空文庫 RoBERTa モデル roberta-base-japanese-luw-upos^[59]を元に、国語研長単位にもとづく品詞付与を考えてみましょう。Google Colaboratory 上の Transformers を用いて、「全学年にわたって小学校の国語の教科書に大量の挿し絵が用いられている」に UPOS を付与するプログラムと、その結果を図 34 に示します。各トークンに系列ラベリングで UPOS を付与しているのですが、Japanese-LUW-Tokenizer が完璧ではなく、やや短めに「全」「学年」「にわたって」「小学校」「の」「国語」「の」「教科書」「に」「大量」「の」「挿し」「絵」「が」「用い」「られて」「いる」とトークナイズしてしまうため、「全」と「挿し」に **B-NOUN**、「学年」と「絵」に **I-NOUN** が付与されています。一方、「大量」に **ADJ** が付与されていますが、これは **NOUN** の方が適切だと考えられます。

続いて、国語研長単位での係り受け解析^[60]も見てみましょう。Google Colaboratory 上の esupar を用いて、「全学年にわたって小学校の国語の教科書に大量の挿し絵が用いられている」に係り受け解析するプログラムと、その結果を図 35 に示します。「小学校の国語の教科書」は、名詞が名詞を修飾していて **nmod** で繋がれていると同時に、「の」が **case** で繋がれています。「大量の挿し絵」も同様のはずですが、「大量」が **NOUN** ではなく **ADJ** になってしまっています。「挿し絵」はガ格であり、**nsubj** で繋がれていると同時に、「が」が **case** で繋がれています。「教科書」は二格であり、**obl** で繋がれていると同時に、「に」が **case** で繋がれています。「全学年」も二格と判断されていて、**obl** で繋

^[58]<https://github.com/KoichiYasuoka/Japanese-LUW-Tokenizer>

^[59]<https://huggingface.co/KoichiYasuoka/roberta-base-japanese-luw-upos>

^[60]安岡孝一: Transformers と国語研長単位による日本語係り受け解析モデルの製作, 情報処理学会研究報告, Vol.2022-CH-128 『人文科学とコンピュータ』, No.7 (2022 年 2 月 19 日), pp.1-8.

```

!pip install deplacy transformers
class SeqL(object):
    def __init__(self,bert):
        from transformers import pipeline
        self.tagger=pipeline(task="ner",model=bert)
    def __call__(self,text):
        w=[(t["start"],t["end"],t["entity"],t["word"]) for t in self.tagger(text)]
        u,z="# text = "+text.replace("\n"," ")+"\n","_"
        for i,(s,e,p,t) in enumerate(w,1):
            m,q=z if i<len(w) and e<w[i][0] else "SpaceAfter=No",p.split("|")
            f=z if p.find("=")<0 else "|".join(t for t in q if t.find("=")>0)
            u+="%t".join([str(i),t.strip(),z,q[0].replace(".", ""),z,f,z,z,z,m])+"\n"
        return u+"\n"
nlp=SeqL("KoichiYasuoka/roberta-base-japanese-luw-upos")
doc=nlp("全学年にわたって小学校の国語の教科書に大量の挿し絵が用いられている")
import deplacy
deplacy.serve(doc,port=None)

```

B-NOUN I-NOUN ADP NOUN ADP NOUN ADP NOUN ADP ADJ ADP
 全 学年 にわたって 小学校 の 国語 の 教科書 に 大量 の

B-NOUN I-NOUN ADP VERB AUX AUX
 挿し 絵 が 用い られ ている

図 34: roberta-base-japanese-luw-upos による UPOS 付与

```

!pip install esupar
import esupar
nlp=esupar.load("KoichiYasuoka/roberta-base-japanese-luw-upos")
doc=nlp("全学年にわたって小学校の国語の教科書に大量の挿し絵が用いられている")
import deplacy
deplacy.serve(doc,port=None)

```

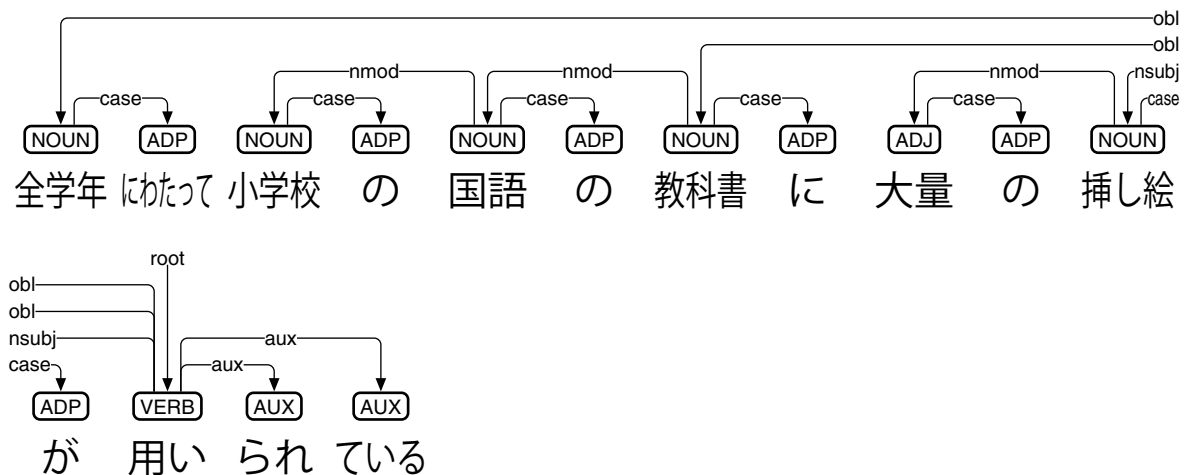


図 35: roberta-base-japanese-luw-upos による日本語係り受け解析 (国語研長単位)

がれていると同時に、「にわたって」が case で繋がられています。

3.4 国語研長単位ミニ DeBERTa モデルの製作

Google Colaboratory 上の esupar で、国語研長単位ミニ DeBERTa モデルを製作してみましょう。図 36 では、図 29～31 で作成した my-dir/deberta-suw を、国語研長単位に使い回す形で、品詞付与・係り受け解析モデル my-dir/deberta-luw-upos を製作しています。国語研短単位の BertJapaneseTokenizer (内部的には fugashi と unidic-lite) を使うことになるので、国語研長単位の UD_Japanese-GSDLUW^[61]に較べてトークンが短くなるのですが、esupar が UPOS 付与に「B-」「I-」を駆使することで、うまく動作するはずです (図 37)。

```
!pip install esupar fugashi unidic-lite spacy-alignments accelerate
url="https://github.com/UniversalDependencies/UD_Japanese-GSDLUW"
!python -m esupar.train my-dir/deberta-suw my-dir/deberta-luw-upos {url}
```

図 36: 国語研長単位向け品詞付与・係り受け解析ミニモデルの製作

```
!pip install esupar fugashi unidic-lite
import esupar
nlp=esupar.load("my-dir/deberta-luw-upos")
doc=nlp("全学年にわたって小学校の国語の教科書に大量の挿し絵が用いられている")
import deplacy
deplacy.serve(doc,port=None)
```

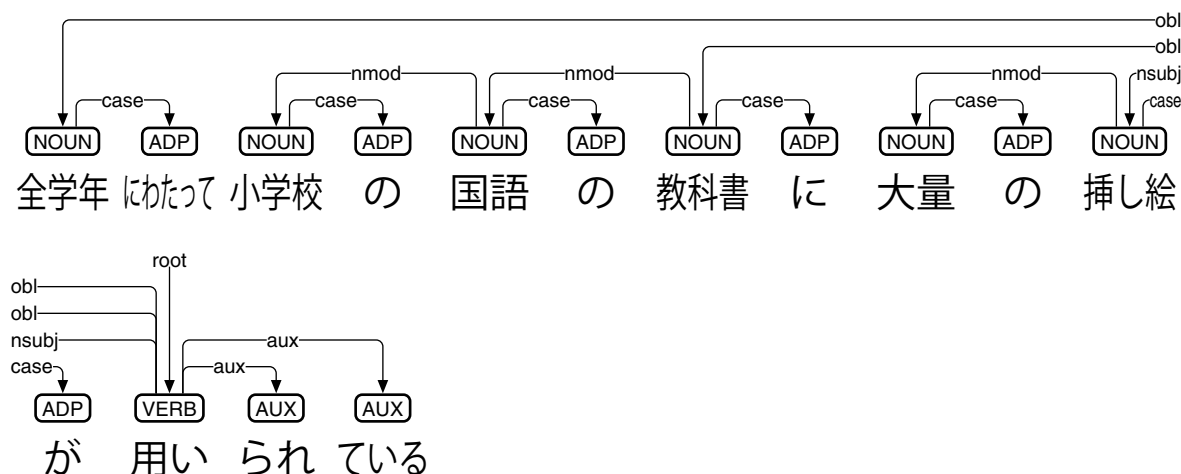


図 37: 国語研長単位向け品詞付与・係り受け解析ミニモデルのテスト

^[61]大村舞, 若狭絢, 浅原正幸: 国語研長単位に基づく日本語 Universal Dependencies, 自然言語処理, Vol.30, No.1 (2023 年 3 月), pp.4-29.

3.5 文節にもとづく係り受け解析

SuPar-UniDic は、国語研短単位にもとづく隣接行列型係り受け解析アルゴリズムを実装しており、その上で、文節にもとづく係り受け解析が可能です。ただし、文節にもとづく係り受けは、慣習的に文頭から文末へとリンクする^[62]ため、矢印の方向が日本語 UD と逆になってしまいます。

Google Colaboratory 上の SuPar-UniDic を用いて、「全学年にわたって小学校の国語の教科書に大量の挿し絵が用いられている」を文節係り受け解析するプログラムと、その結果(可視化は Graphviz^[63]を使用)を図 38 に示します。「全学年にわたって」が1文節なのか2文節なのか議論の余地があるのですが、文節にもとづく係り受け解析は UD と全く異なっていることから、ここでは参考にとめておきます。

```
!pip install suparunidic
import suparunidic
nlp=suparunidic.load("gendai")
doc=nlp("全学年にわたって小学校の国語の教科書に大量の挿し絵が用いられている")
d=suparunidic.bunsetsu_spans(doc)
from deplacy.deprelja import deprelja
g="digraph{";".join(['x{}[label="{}"]'.format(b.start,b.text) for b in d]+
['x{}->x{}[label="{}",fontsize=9]'.format(suparunidic.bunsetsu_span(t).start,
b.start,deprelja[t.dep_]) for b in d for t in b.lefts])+"}"
import graphviz
graphviz.Source(g)
```

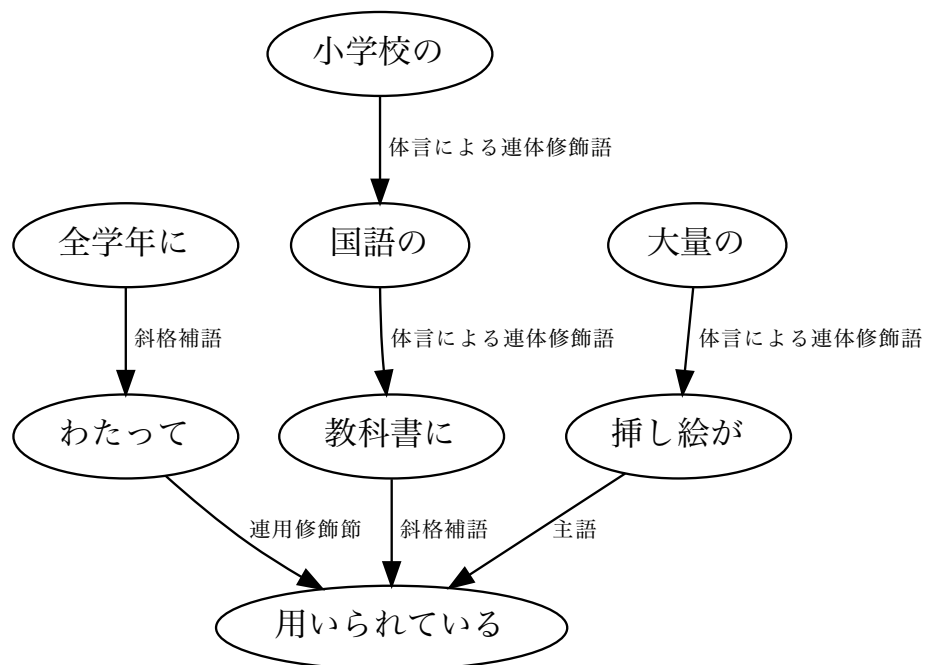


図 38: 文節にもとづく日本語係り受け解析

^[62]吉田将: 二文節間の係り受けを基礎とした日本語文の構文分析, 電子通信学会論文誌, Vol.55-D, No.4 (1972 年 4 月), pp.238-244.

^[63]<https://graphviz.readthedocs.io>

4 英語編

4.1 英語 UD における品詞付与と係り受け解析

Trankit^[64]は、XLM-RoBERTa^[65]をベースに、隣接行列型係り受け解析アルゴリズムを、様々な言語向けに実装しています。各言語向けに品詞付与も実装していて、複数の言語を同時に扱うことも可能です。Google Colaboratory 上の Trankit を用いて、英文「It didn't take long to repair my car」を品詞付与・係り受け解析するプログラムと、その結果を図 39 に示します。解析結果を見ていくことにしましょう。

```
!pip install git+https://github.com/KoichiYasuoka/trankit deplacy
import trankit
nlp=trankit.Pipeline("english")
doc=nlp("It didn't take long to repair my car")
import deplacy
deplacy.serve(doc,port=None)
```

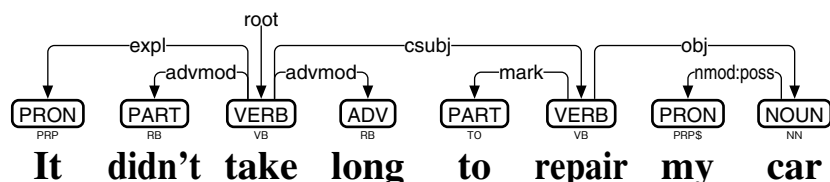


図 39: Trankit による英語の係り受け解析

動詞「take」の主語は「It」ですが、「It」は形式主語で、実際的主語は「to repair my car」という節 (clause) です。これを反映して、「It」が **expl** で、「repair」が **csubj** で繋がられています。「didn't」と「long」は「take」を連用修飾していて、いずれも **advmod** で繋がられています。「to repair my car」においては、「to」が節全体に前置されていて、**mark** で繋がられています。「car」は「repair」の目的語で、**obj** で繋がられています。「my」は「car」を修飾しており、本来は **det** が適切だと考えられるのですが、英語 UD では **nmod:poss** (所有に関する連体修飾) を導入しています。なお、Trankit は「didn't」を1語だとみなしていますが、英語 UD では「did」「n't」の2語に分けるべきであり、単語切りが不十分です。

UDPipe 2 WebAPI^[66]は、UDPipe 2 の係り受け解析モデル^[67]を、WebAPI の形で公開しています。内部的には、隣接行列型係り受け解析アルゴリズムを使っています。Google Colaboratory から UDPipe 2 WebAPI の english-ewt にアクセスするプログラムと、その結

^[64]Minh Van Nguyen, Viet Dac Lai, Amir Pouran Ben Veyseh and Thien Huu Nguyen: Trankit: A Light-Weight Transformer-based Toolkit for Multilingual Natural Language Processing, EACL 2021: 16th Conference of the European Chapter of the Association for Computational Linguistics (April 2021), System Demonstrations, pp.80-90.

^[65]Alexis Conneau, Kartikay Khandelwal, Naman Goyal, Vishrav Chaudhary, Guillaume Wenzek, Francisco Guzmán, Edouard Grave, Myle Ott, Luke Zettlemoyer, Veselin Stoyanov: Unsupervised Cross-lingual Representation Learning at Scale, Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics (July 2020), pp.8440-8451.

^[66]<https://lindat.mff.cuni.cz/services/udpipe/>

^[67]Milan Straka: UDPipe 2.0 Prototype at CoNLL 2018 UD Shared Task, Proceedings of the CoNLL 2018 Shared Task (October 2018), pp.197-207.

```

!pip install deplacy
class UDPipe2WebAPI(object):
    def __init__(self, lang):
        self.url="https://lindat.mff.cuni.cz/services/udpipe/api/process"
        self.params={"model":lang,"tokenizer":"","tagger":"","parser":""}
    def __call__(self, text):
        import urllib.parse,urllib.request,json
        self.params["data"]=urllib.parse.quote(text)
        u=self.url+"?"+"&".join(k+"="+v for k,v in self.params.items())
        with urllib.request.urlopen(u) as r:
            return json.loads(r.read())["result"]
nlp=UDPipe2WebAPI("english-ewt")
doc=nlp("It didn't take long to repair my car")
import deplacy
deplacy.serve(doc,port=None)

```

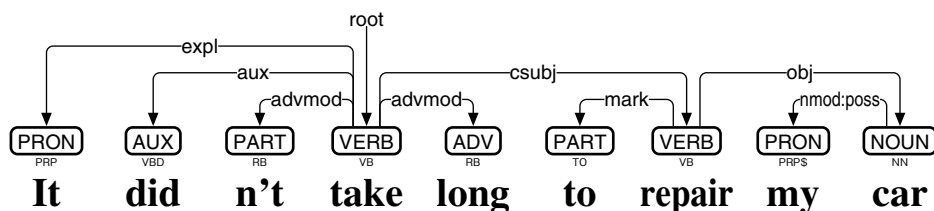


図 40: UDPipe 2 WebAPI による英語の係り受け解析

果を図 40 に示します。「It didn't take long to repair my car」の解析において、「did」「n't」が2語に分かれていて、それぞれ **aux** と **advmod** で繋がられています。

なお、英語 UD の係り受けタグには、表 4 と **nmod:poss** に加え、斜格補語による連用修飾を表す **obl:npm**、時に関する連体修飾を表す **nmod:tmod**、時に関する斜格補語を表す **obl:tmod**、受動態の主語を表す **nsubj:pass**、受動態の節主語を表す **csubj:pass**、関係代名詞による連体修飾節を表す **acl:relcl**、限定詞の前に置かれる形容詞を表す **det:predet**、接続における前置副詞を表す **cc:preconj**、句動詞を表す **compound:p** も用いられます^[68]。

4.2 英語ミニ DeBERTa モデルの製作

Google Colaboratory 上の `esupar`^[47] で、英語ミニ DeBERTa モデルを製作してみましょう。手順を図 41～44 に示します。

最初に、訓練のための文章を準備しましょう (図 41)。train.txt の文章の材料として、ここでは UD_English-EWT^[69] と WikiText-2^[70] を元にしてしていますが、好みに応じ、他の英文を増量するのもいいでしょう。なお、UD_English-EWT の各 CoNLL-U ファイルは、品詞付与・係り受け解析の学習にも用います。

続けて、英語向けトークナイザを、my-dir/tokenizer-en に作成しましょう。図 42 では、Transformers^[26] の BertTokenizerFast を使って、英語向けトークナイザを作成しています。

^[68]<https://universaldependencies.org/en/#syntax>

^[69]https://github.com/UniversalDependencies/UD_English-EWT

^[70]Stephen Merity, Caiming Xiong, James Bradbury and Richard Socher: Pointer Sentinel Mixture Models, ICLR 2017: 5th International Conference on Learning Representations, Wednesday Afternoon Poster Session C27 (April 26, 2017).

```

!pip install esupar spacy-alignments accelerate datasets
import os
from datasets import load_dataset
url="https://github.com/UniversalDependencies/UD_English-EWT"
d=os.path.basename(url)
!test -d {d} || git clone --depth=1 {url}
!for F in train dev test ; do cp {d}/*-$$F.conllu $$F.conllu ; done
with open("train.txt","w",encoding="utf-8") as w:
    for k in load_dataset("Salesforce/wikitext","wikitext-2-raw-v1")["train"]:
        print(k["text"],file=w)
!sed -n "s/^# text = //p" *.conllu >> train.txt

```

図 41: 英語ミニモデル製作のための準備

```

from transformers import BertTokenizerFast
from tokenizers import BertWordPieceTokenizer
s=[" [CLS] ", " [PAD] ", " [SEP] ", " [UNK] ", " [MASK] "]
wpt=BertWordPieceTokenizer()
wpt.train(files=["train.txt"],vocab_size=8000,special_tokens=s)
wpt.save_model(".")
tkz=BertTokenizerFast(vocab_file="vocab.txt",never_split=s)
tkz.save_pretrained("my-dir/tokenizer-en")

```

図 42: 英語向けトークナイザの作成

```

from transformers import (AutoTokenizer,DebertaV2Config,DebertaV2ForMaskedLM,
    DataCollatorForLanguageModeling,TrainingArguments,Trainer)
tkz=AutoTokenizer.from_pretrained("my-dir/tokenizer-en",model_max_length=128)
t=tkz.convert_tokens_to_ids([" [CLS] ", " [PAD] ", " [SEP] ", " [UNK] ", " [MASK] "])
cfg=DebertaV2Config(hidden_size=256,num_hidden_layers=12,num_attention_heads=4,
    intermediate_size=768,relative_attention=True,position_biased_input=False,
    pos_att_type=["p2c","c2p"],max_position_embeddings=tkz.model_max_length,
    vocab_size=len(tkz),tokenizer_class=type(tkz).__name__,
    bos_token_id=t[0],pad_token_id=t[1],eos_token_id=t[2])
arg=TrainingArguments(report_to="none",output_dir="/tmp",save_total_limit=2,
    num_train_epochs=3,per_device_train_batch_size=64)
class ReadLineDS(object):
    def __init__(self,file,tokenizer):
        self.tokenizer=tokenizer
        with open(file,"r",encoding="utf-8") as r:
            self.lines=[s.strip() for s in r if s.strip()!=""]
    __len__=lambda self:len(self.lines)
    __getitem__=lambda self,i:self.tokenizer(self.lines[i],truncation=True,
        add_special_tokens=True,max_length=self.tokenizer.model_max_length-2)
trn=Trainer(args=arg,data_collator=DataCollatorForLanguageModeling(tkz),
    model=DebertaV2ForMaskedLM(cfg),train_dataset=ReadLineDS("train.txt",tkz))
trn.train()
trn.save_model("my-dir/deberta-en")
tkz.save_pretrained("my-dir/deberta-en")

```

図 43: 英語ミニ DeBERTa モデルの作成

語彙数 (vocab_size) を 8000 に絞っていますが、もう少し大きくした方がいいかもしれません。

次に train.txt の文章から、DeBERTa モデルを my-dir/deberta-en に作成しましょう。図 43 では、入出力幅 128 トークン・入出力ベクトル 256 次元・深さ 12 層・アテンションヘッド 4 個・中間ベクトル 768 次元という、やや小さめのモデルにしており、CPU なら 2 時間程度、GPU なら 20 分程度で作成できます。

最後に、my-dir/deberta-en と UD_English-EWT から、品詞付与・係り受け解析モデル my-dir/deberta-en-upos を製作しましょう (図 44)。うまく my-dir/deberta-en-upos が製作できたら、esupar でテストしてみましょう (図 45)。

```
!python -m esupar.train my-dir/deberta-en my-dir/deberta-en-upos .
```

図 44: 英語向け品詞付与・係り受け解析ミニモデルの製作

```
!pip install esupar
import esupar
nlp=esupar.load("my-dir/deberta-en-upos")
doc=nlp("It didn't take long to repair my car")
import deplacy
deplacy.serve(doc,port=None)
```

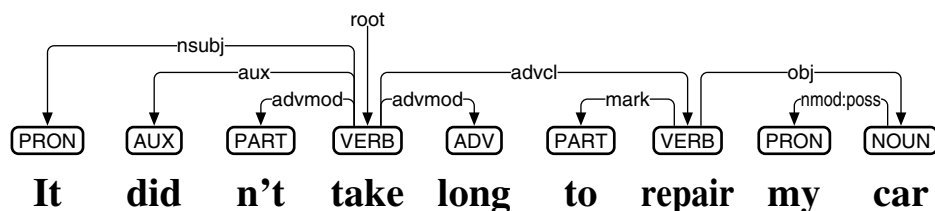


図 45: 英語向け品詞付与・係り受け解析ミニモデルのテスト

4.3 英語 UD による構成素解析

Stanza^[71]は、隣接行列型係り受け解析アルゴリズムを実装しており、それに加えて、状態遷移型構成素解析アルゴリズム^[72]も実装しています。Google Colaboratory 上の Stanza を用いて、「It didn't take long to repair my car」を構成素解析するプログラムと、その結果 (可視化は Graphviz^[63]を使用) を図 46 に示します。構成素解析は、UD とは全く異なっていることから、ここでは参考に留めておきます。

^[71]Peng Qi, Yuhao Zhang, Yuhui Zhang, Jason Bolton, Christopher D. Manning: Stanza: A Python Natural Language Processing Toolkit for Many Human Languages, 58th Annual Meeting of the Association for Computational Linguistics: Proceedings of the System Demonstration (July 2020), pp.101-108.

^[72]Jiangming Liu and Yue Zhang: In-Order Transition-based Constituent Parsing, Transactions of the Association for Computational Linguistics, Vol.5, pp.413-424 (November 2017).

```

!pip install stanza
import stanza
nlp=stanza.Pipeline("en")
doc=nlp("It didn't take long to repair my car")
tree2dot=lambda x,y="x":('graph{{node[shape="plaintext"];{}}}' if len(y)<2
    else y[0:-1]+"--"+y+";{}}").format(";\\n".join([y+'[label="'+x.label+'"]']+
    [tree2dot(c,y+chr(i+97)) for i,c in enumerate(x.children)]))
import graphviz
graphviz.Source(tree2dot(doc.sentences[0].constituency))

```

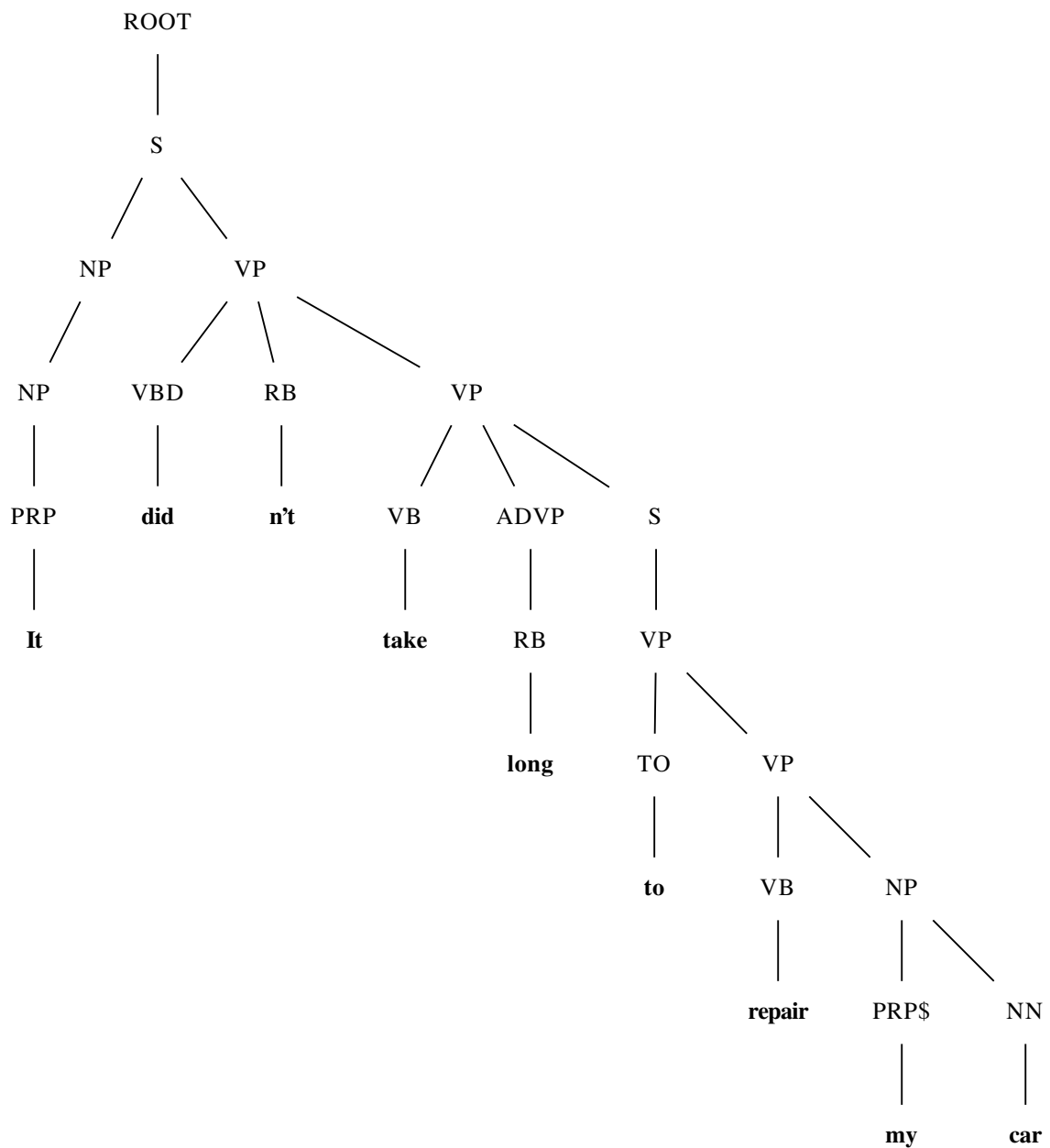


図 46: Stanza による英語の構成素解析

5 フランス語編

5.1 フランス語UDにおける品詞付与と係り受け解析

spacy-transformers^[73]は、事前学習モデルによる状態遷移型係り受け解析アルゴリズムを実装しており、フランス語向けには、CamemBERT^[74]を元にした fr_dep_news_trf を公開^[75]しています。Google Colaboratory 上の spacy-transformers を用いて、フランス語の例文「Il ne peut pas prêter attention aux articles contractés」を品詞付与・係り受け解析するプログラムと、その結果を図47に示します。解析結果を見ていくことにしましょう。

```
!pip install spacy-transformers deplacy
!python -m spacy download fr_dep_news_trf
import spacy
nlp=spacy.load("fr_dep_news_trf")
doc=nlp("Il ne peut pas prêter attention aux articles contractés")
import deplacy
deplacy.serve(doc,port=None)
```

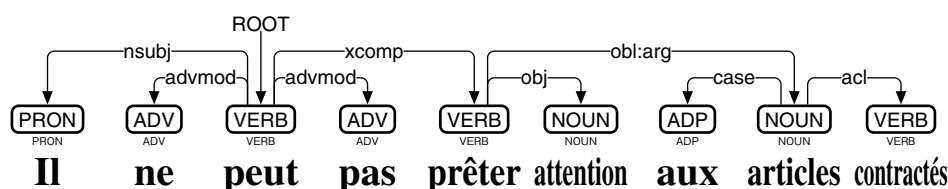


図 47: spacy-transformers によるフランス語の係り受け解析

動詞「peut」の主語「Il」は、同時に「prêter」の主語でもあることから、「Il」が nsubj で、「prêter」が xcomp で繋がれています。「ne」「pas」は「peut」を否定しているとみなして、いずれも advmod で繋がれています。「attention」は「prêter」の目的語であり、obj で繋がれています。斜格補語「articles」については、説明が必要です。フランス語 UD では、obl を obl:agent・obl:arg・obl:mod の3つに分離しており、それぞれ動作主 (agent)・項 (argument)・修飾子 (modificateur) を表しています。「articles」は obl:arg で繋がれていることから、「prêter」の項 (ただし動作主ではない) だとみなされています。「articles」に対し「aux」は前置詞、「contractés」は連体修飾節だとみなされており、それぞれ case・acl で繋がれています。なお、spacy-transformers は「aux」を1語だとみなしていますが、フランス語 UD では「à」「les」の2語に分離した上で扱うことになっており、縮合語の解析が不十分です。

Trankit^[64]は、XLM-RoBERTa をベースに、隣接行列型係り受け解析アルゴリズムを、様々な言語向けに実装しています。Google Colaboratory 上の Trankit を用いて「Il ne peut pas prêter attention aux articles contractés」を品詞付与・係り受け解析するプログラムと、その結果を図48に示します。「aux」が内部的には2語として扱われており、「à」が「articles」の前置詞、「les」が限定詞 (冠詞) として、それぞれ case・det で繋がれています。

^[73]<https://github.com/explosion/spacy-transformers>

^[74]Louis Martin, Benjamin Muller, Pedro Javier Ortiz Suárez, Yoann Dupont, Laurent Romary, Éric Villemonte de la Clergerie, Djamé Seddah, Benoît Sagot: CamemBERT: a Tasty French Language Model, Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics (July 2020), pp.7203-7219.

^[75]https://huggingface.co/spacy/fr_dep_news_trf

```
!pip install git+https://github.com/KoichiYasuoka/trankit deplacy
import trankit
nlp=trankit.Pipeline("french")
doc=nlp("Il ne peut pas prêter attention aux articles contractés")
import deplacy
deplacy.serve(doc,port=None)
```

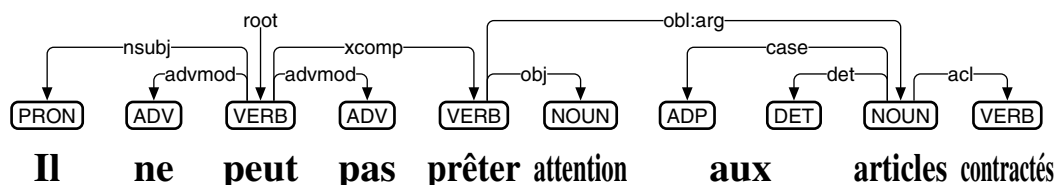


図 48: Trankit によるフランス語の係り受け解析

なお、フランス語 UD の係り受けタグには、表 4 と `obl:agent`・`obl:arg`・`obl:mod`に加え、数多くの拡張が用いられている^[76]のですが、拡張どうしの中で混乱が起こっており、事態の收拾が望まれます。

5.2 フランス語ミニ DeBERTa モデルの製作

Google Colaboratory 上の `esupar`^[47]で、フランス語ミニ DeBERTa モデルを製作してみましょう。手順を図 49～52 に示します。

最初に、訓練のための文章を準備しましょう (図 49)。`train.txt` の文章の材料として、ここでは UD_French-Sequoia^[77]と UD_French-GSD^[78]を元にしていますが、好みに応じ、フランス語の他の文章を増量するのもいいでしょう。なお、各 CoNLL-U ファイルは、品詞付与・係り受け解析の学習にも用います。

続けて、フランス語向けトークナイザを、`my-dir/tokenizer-fr`に作成しましょう。図 50 では、Transformers^[26]の `BertTokenizerFast` を使って、フランス語向けトークナイザを作成しています。

次に `train.txt` の文章から、DeBERTa モデルを `my-dir/deberta-fr`に作成しましょう。図 51 では、入出力幅 128 トークン・入出力ベクトル 256 次元・深さ 12 層・アテンションヘッ

```
!pip install esupar spacy-alignments accelerate
import os
url="https://github.com/UniversalDependencies/UD_French-Sequoia"
!test -d {os.path.basename(url)} || git clone --depth=1 {url}
url="https://github.com/UniversalDependencies/UD_French-GSD"
!test -d {os.path.basename(url)} || git clone --depth=1 {url}
!for F in train dev test ; do cat UD_French-*/*-$$F.conllu > $$F.conllu ; done
!sed -n "s/^# text = //p" *.conllu > train.txt
```

図 49: フランス語ミニモデル製作のための準備

^[76]<https://universaldependencies.org/fr/dep>

^[77]https://github.com/UniversalDependencies/UD_French-Sequoia

^[78]https://github.com/UniversalDependencies/UD_French-GSD

```

from transformers import BertTokenizerFast
from tokenizers import BertWordPieceTokenizer
s=[" [CLS]", " [PAD]", " [SEP]", " [UNK]", " [MASK]"]
wpt=BertWordPieceTokenizer(strip_accents=False)
wpt.train(files=["train.txt"],vocab_size=30000,special_tokens=s)
wpt.save_model(".")
tkz=BertTokenizerFast(vocab_file="vocab.txt",strip_accents=False,never_split=s)
tkz.save_pretrained("my-dir/tokenizer-fr")

```

図 50: フランス語向けトークナイザの作成

ド 4 個・中間ベクトル 768 次元という、やや小さめのモデルにしており、CPU なら 1 時間程度、GPU なら 10 分程度で作成できます。

```

from transformers import (AutoTokenizer,DebertaV2Config,DebertaV2ForMaskedLM,
    DataCollatorForLanguageModeling,TrainingArguments,Trainer)
tkz=AutoTokenizer.from_pretrained("my-dir/tokenizer-fr",model_max_length=128)
t=tkz.convert_tokens_to_ids([" [CLS]", " [PAD]", " [SEP]", " [UNK]", " [MASK]"])
cfg=DebertaV2Config(hidden_size=256,num_hidden_layers=12,num_attention_heads=4,
    intermediate_size=768,relative_attention=True,position_biased_input=False,
    pos_att_type=["p2c","c2p"],max_position_embeddings=tkz.model_max_length,
    vocab_size=len(tkz),tokenizer_class=type(tkz).__name__,
    bos_token_id=t[0],pad_token_id=t[1],eos_token_id=t[2])
arg=TrainingArguments(report_to="none",output_dir="/tmp",save_total_limit=2,
    num_train_epochs=3,per_device_train_batch_size=64)
class ReadLineDS(object):
    def __init__(self,file,tokenizer):
        self.tokenizer=tokenizer
        with open(file,"r",encoding="utf-8") as r:
            self.lines=[s.strip() for s in r if s.strip()!=""]
    __len__=lambda self:len(self.lines)
    __getitem__=lambda self,i:self.tokenizer(self.lines[i],truncation=True,
        add_special_tokens=True,max_length=self.tokenizer.model_max_length-2)
trn=Trainer(args=arg,data_collator=DataCollatorForLanguageModeling(tkz),
    model=DebertaV2ForMaskedLM(cfg),train_dataset=ReadLineDS("train.txt",tkz))
trn.train()
trn.save_model("my-dir/deberta-fr")
tkz.save_pretrained("my-dir/deberta-fr")

```

図 51: フランス語ミニ DeBERTa モデルの作成

最後に、my-dir/deberta-fr と UD_French-Sequoia・UD_French-GSD から、品詞付与・係り受け解析モデル my-dir/deberta-fr-upos を製作しましょう (図 52)。ただ、隣接行列型アルゴリズムの学習は収束が遅く、GPU でも 3 時間程度かかってしまいます。

うまく my-dir/deberta-fr-upos が製作できたら、esupar でテストしてみましょう (図 53)。esupar は縮合語を UPOS 付与においてサポート (図 54) しており、係り受け解析では「aux」を 2 語として処理します。

```
!python -m esupar.train my-dir/deberta-fr my-dir/deberta-fr-upos .
```

図 52: フランス語向け品詞付与・係り受け解析ミニモデルの製作

```
!pip install esupar
import esupar
nlp=esupar.load("my-dir/deberta-fr-upos")
doc=nlp("Il ne peut pas prêter attention aux articles contractés")
import deplacy
deplacy.serve(doc,port=None)
```

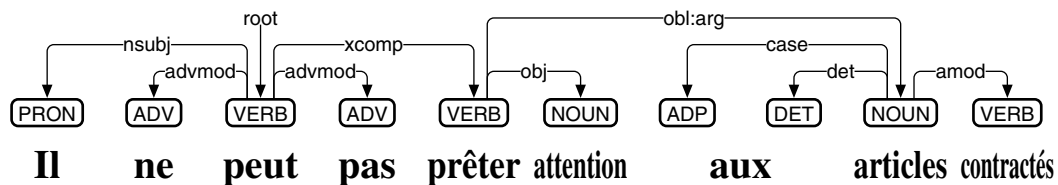


図 53: フランス語向け品詞付与・係り受け解析ミニモデルのテスト

```
!pip install deplacy transformers
class SeqL(object):
    def __init__(self,bert):
        from transformers import pipeline
        self.tagger=pipeline(task="ner",model=bert)
    def __call__(self,text):
        w=[(t["start"],t["end"],t["entity"],t["word"]) for t in self.tagger(text)]
        u,z="# text = "+text.replace("\n"," ")+"\n","_"
        for i,(s,e,p,t) in enumerate(w,1):
            m,q=z if i<len(w) and e<w[i][0] else "SpaceAfter=No",p.split("|")
            f=z if p.find("=")<0 else "|".join(t for t in q if t.find("=")>0)
            u+="\t".join([str(i),t.strip(),z,q[0].replace(".",","),z,f,z,z,z,m])+"\n"
        return u+"\n"
nlp=SeqL("my-dir/deberta-fr-upos")
doc=nlp("Il ne peut pas prêter attention aux articles contractés")
import deplacy
deplacy.serve(doc,port=None)
```

PRON ADV VERB ADV VERB NOUN ADP+DET NOUN B-VERB I-VERB
 Il ne peut pas prêter attention aux articles contract és

図 54: フランス語向け品詞付与・係り受け解析ミニモデルでの UPOS 付与

6 ドイツ語編

6.1 ドイツ語 UD における品詞付与と係り受け解析

modernbert-german-134m-ud-embeds^[79]は、隣接行列型係り受け解析アルゴリズムを、事前学習モデルで実装しています。Google Colaboratory 上の Transformers を用いて、ドイツ語の例文「Im Krankenzimmer sah er krank aus aber wollte gut aussehen」を modernbert-german-134m-ud-embeds で品詞付与・係り受け解析するプログラムと、その結果を図 55 に示します。

```
!pip install deplacy transformers
from transformers import pipeline
mdl="KoichiYasuoka/modernbert-german-134m-ud-embeds"
nlp=pipeline("universal-dependencies",mdl,trust_remote_code=True)
doc=nlp("Im Krankenzimmer sah er krank aus aber wollte gut aussehen")
import deplacy
deplacy.serve(doc,port=None)
```

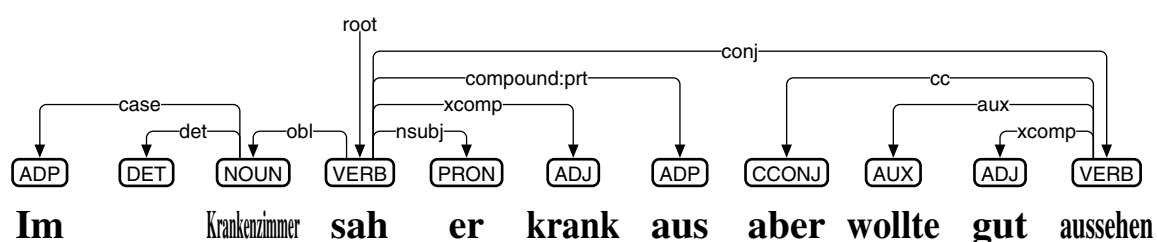


図 55: modernbert-german-134m-ud-embeds によるドイツ語の係り受け解析

ドイツ語 UD は、単語の扱いが非常に独特です。図 55 の「Im Krankenzimmer」では、縮合語「im」を前置詞「in」と冠詞「dem」の2語に分離して扱う一方、複合語である「Krankenzimmer」は1語として扱っています。この結果、「Krankenzimmer」から前置詞へ **case** が、冠詞へ **det** が繋がれています。「aussehen」など分離動詞は、分離していない場合は1語として扱い、分離している場合は2語として扱うと同時に **compound:prt** で繋ぐ、というルールが採用されています。結果として「sah」から「aus」へ **compound:prt** が繋がれています。また「sah」からは、主語「er」が **nsubj** で、補語「krank」が **xcomp** で、斜格補語「Krankenzimmer」が **obl** で、それぞれ繋がれています。「sah aus」と「aussehen」は並置されていて、**conj** で繋がれています。「aussehen」からは、接続詞「aber」が **cc** で、助動詞「wollte」が **aux** で、補語「gut」が **xcomp** で繋がれています。

なお、ドイツ語 UD の係り受けタグには、表 4 と **compound:prt** に加え、受動態を表す **nsubj:pass**・**csbj:pass**・**aux:pass**、斜格補語の動作主を表す **obl:agent**、項(動作主以外)の斜格補語を表す **obl:arg**、所有を表す **det:pos**・**nmod:poss**、再帰動詞における再帰代名詞を表す **expl:pv** も用いられます^[80]。

^[79]<https://huggingface.co/KoichiYasuoka/modernbert-german-134m-ud-embeds>

^[80]<https://universaldependencies.org/de/#relations-overview>

6.2 ドイツ語ミニ DeBERTa モデルの製作

Google Colaboratory 上の esupar で、ドイツ語ミニ DeBERTa モデルを製作してみましょう。手順を図 56～59 に示します。

最初に、訓練のための文章を準備しましょう (図 56)。train.txt の文章の材料として、ここでは UD_German-HDT^[81]を元にしていますが、好みに応じ、ドイツ語の他の文章を増量するのもいいでしょう。なお、UD_German-HDT の CoNLL-U ファイルは、品詞付与・係り受け解析の学習にも用います。

```
!pip install esupar spacy-alignments accelerate
import os
url="https://github.com/UniversalDependencies/UD_German-HDT"
d=os.path.basename(url)
!test -d {d} || git clone --depth=1 {url}
!for F in train dev test ; do cat {d}/*-$$F*.conllu > $$F.conllu ; done
!sed -n "s/^# text = //p" *.conllu > train.txt
```

図 56: ドイツ語ミニモデル製作のための準備

続けて、ドイツ語向けトークナイザを、my-dir/tokenizer-de に作成しましょう。図 57 では、Transformers^[26]の BertTokenizerFast を使って、ドイツ語向けトークナイザを作成しています。

```
from transformers import BertTokenizerFast
from tokenizers import BertWordPieceTokenizer
s=["[CLS]","[PAD]","[SEP]","[UNK]","[MASK]"]
wpt=BertWordPieceTokenizer(strip_accents=False)
wpt.train(files=["train.txt"],vocab_size=30000,special_tokens=s)
wpt.save_model(".")
tkz=BertTokenizerFast(vocab_file="vocab.txt",strip_accents=False,never_split=s)
tkz.save_pretrained("my-dir/tokenizer-de")
```

図 57: ドイツ語向けトークナイザの作成

次に train.txt の文章から、DeBERTa モデルを my-dir/deberta-de に作成しましょう。図 58 では、入出力幅 128 トークン・入出力ベクトル 256 次元・深さ 12 層・アテンションヘッド 4 個・中間ベクトル 768 次元という、やや小さめのモデルにしていますが、GPU でも 1 時間程度かかってしまいます。

最後に、my-dir/deberta-de と UD_German-HDT から、品詞付与・係り受け解析モデル my-dir/deberta-de-upos を製作しましょう (図 59)。うまく my-dir/deberta-de-upos が製作できたら、esupar でテストしてみましょう (図 60)。esupar は縮合語を UPOS 付与においてサポート (図 61) しており、係り受け解析では「Im」を 2 語として処理します。

^[81]https://github.com/UniversalDependencies/UD_German-HDT

```

from transformers import (AutoTokenizer,DebertaV2Config,DebertaV2ForMaskedLM,
    DataCollatorForLanguageModeling,TrainingArguments,Trainer)
tkz=AutoTokenizer.from_pretrained("my-dir/tokenizer-de",model_max_length=128)
t=tkz.convert_tokens_to_ids(["[CLS]","[PAD]","[SEP]","[UNK]","[MASK]"])
cfg=DebertaV2Config(hidden_size=256,num_hidden_layers=12,num_attention_heads=4,
    intermediate_size=768,relative_attention=True,position_biased_input=False,
    pos_att_type=["p2c","c2p"],max_position_embeddings=tkz.model_max_length,
    vocab_size=len(tkz),tokenizer_class=type(tkz).__name__,
    bos_token_id=t[0],pad_token_id=t[1],eos_token_id=t[2])
arg=TrainingArguments(report_to="none",output_dir="/tmp",save_total_limit=2,
    num_train_epochs=3,per_device_train_batch_size=64)
class ReadLineDS(object):
    def __init__(self,file,tokenizer):
        self.tokenizer=tokenizer
        with open(file,"r",encoding="utf-8") as r:
            self.lines=[s.strip() for s in r if s.strip()!=""]
    __len__=lambda self:len(self.lines)
    __getitem__=lambda self,i:self.tokenizer(self.lines[i],truncation=True,
        add_special_tokens=True,max_length=self.tokenizer.model_max_length-2)
trn=Trainer(args=arg,data_collator=DataCollatorForLanguageModeling(tkz),
    model=DebertaV2ForMaskedLM(cfg),train_dataset=ReadLineDS("train.txt",tkz))
trn.train()
trn.save_model("my-dir/deberta-de")
tkz.save_pretrained("my-dir/deberta-de")

```

図 58: ドイツ語ミニ DeBERTa モデルの作成

```
!python -m esupar.train my-dir/deberta-de my-dir/deberta-de-upos .
```

図 59: ドイツ語向け品詞付与・係り受け解析ミニモデルの製作

```

!pip install esupar
import esupar
nlp=esupar.load("my-dir/deberta-de-upos")
doc=nlp("Im Krankenzimmer sah er krank aus aber wollte gut aussehen")
import deplacy
deplacy.serve(doc,port=None)

```

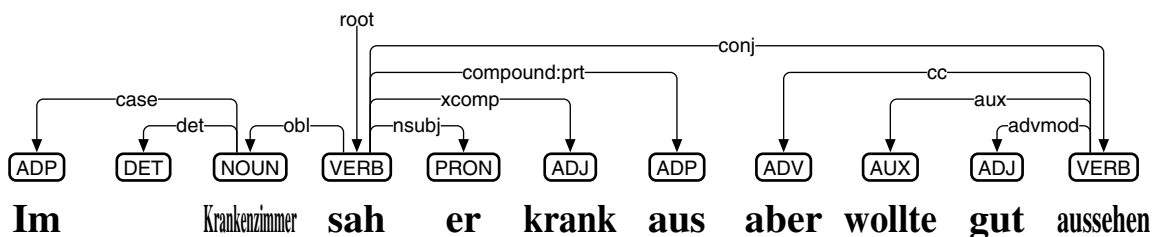


図 60: ドイツ語向け品詞付与・係り受け解析ミニモデルのテスト

```

!pip install deplacy transformers
class SeqL(object):
    def __init__(self,bert):
        from transformers import pipeline
        self.tagger=pipeline(task="ner",model=bert)
    def __call__(self,text):
        w=[(t["start"],t["end"],t["entity"],t["word"]) for t in self.tagger(text)]
        u,z="# text = "+text.replace("\n"," ")+"\n","_"
        for i,(s,e,p,t) in enumerate(w,1):
            m,q=z if i<len(w) and e<w[i][0] else "SpaceAfter=No",p.split("|")
            f=z if p.find("=")<0 else "|".join(t for t in q if t.find("=")>0)
            u+="\t".join([str(i),t.strip(),z,q[0].replace(".", ""),z,f,z,z,z,m])+"\n"
        return u+"\n"
nlp=SeqL("my-dir/deberta-de-upos")
doc=nlp("Im Krankenzimmer sah er krank aus aber wollte gut aussehen")
import deplacy
deplacy.serve(doc,port=None)

```

ADP+DET B-NOUN I-NOUN VERB PRON ADJ ADP ADV AUX ADJ VERB
Im Krankenzimmer sah er krank aus aber wollte gut aussehen

図 61: ドイツ語向け品詞付与・係り受け解析ミニモデルでの UPOS 付与

7 オランダ語編

7.1 オランダ語UDにおける品詞付与と係り受け解析

Trankit^[64]は、XLM-RoBERTa をベースに、隣接行列型係り受け解析アルゴリズムを、様々な言語向けに実装しています。Google Colaboratory 上の Trankit を用いて、オランダ語の例文「In de ziekenkamer zag hij ziek uit maar wilde goed uitzien」を品詞付与・係り受け解析するプログラムと、その結果を図 62 に示します。解析結果を見ていくことにしましょう。

```
!pip install git+https://github.com/KoichiYasuoka/trankit deplacy
import trankit
nlp=trankit.Pipeline("dutch")
doc=nlp("In de ziekenkamer zag hij ziek uit maar wilde goed uitzien")
import deplacy
deplacy.serve(doc,port=None)
```

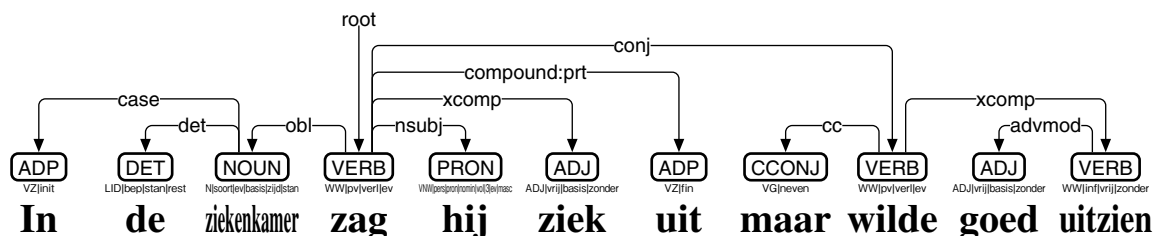


図 62: Trankit によるオランダ語の係り受け解析

「uitzien」など分離動詞は、分離していない場合は1語として扱い、分離している場合は2語として扱うと同時に **compound:prt** で繋ぐ、というルールが採用されています。結果として「zag」から「uit」へ **compound:prt** が繋がられています。また「zag」からは、主語「hij」が **nsubj** で、補語「ziek」が **xcomp** で、斜格補語「zielenkamer」が **obl** で、それぞれ繋がられています。「zielenkamer」からは、前置詞「In」が **case** で、冠詞「de」が **det** で繋がられています。「zag uit」と「wilde uitzien」は並置されていて、**conj** で繋がられています。「wilde」には接続詞「maar」が **cc** で、「uitzien」が **xcomp** で繋がられています。「uitzien」と「goed」は **advmod** で繋がっていますが、これは **xcomp** の方が適切だと考えられます。

なお、オランダ語UDの係り受けタグには、表4と **compound:prt** に加え、受動態を表す **nsubj:pass** と **aux:pass**、受動態の動作主を表す **obl:agent**、所有を表す **nmod:pos**、関係代名詞による連体修飾節を表す **acl:relcl**、再帰動詞における再帰代名詞を表す **expl:pv** も用いられます^[82]。

7.2 オランダ語ミニ DeBERTa モデルの製作

Google Colaboratory 上の esupar^[47]で、オランダ語ミニ DeBERTa モデルを製作してみましょう。手順を図 63～66 に示します。

^[82]<https://universaldependencies.org/nl/#syntax>

最初に、訓練のための文章を準備しましょう (図 63)。train.txt の文章の材料として、ここでは UD_Dutch-Alpino^[83] と UD_Dutch-LassySmall^[84] を元にしてありますが、好みに応じ、オランダ語の他の文章を増量するのもいいでしょう。なお、各 CoNLL-U ファイルは、品詞付与・係り受け解析の学習にも用います。

```
!pip install esupar accelerate
import os
url="https://github.com/UniversalDependencies/UD_Dutch-Alpino"
!test -d {os.path.basename(url)} || git clone --depth=1 {url}
url="https://github.com/UniversalDependencies/UD_Dutch-LassySmall"
!test -d {os.path.basename(url)} || git clone --depth=1 {url}
!for F in train dev test ; do cat UD_Dutch-*/-$$F.conllu > $$F.conllu ; done
!sed -n "s/^# text = //p" *.conllu > train.txt
```

図 63: オランダ語ミニモデル製作のための準備

続けて、オランダ語向けトークナイザを、my-dir/tokenizer-nl に作成しましょう。図 64 では、Transformers^[26] の BertTokenizerFast を使って、オランダ語向けトークナイザを作成しています。語彙数 (vocab_size) を 8000 に絞っていますが、もう少し大きくした方がいいかもしれません。

```
from transformers import BertTokenizerFast
from tokenizers import BertWordPieceTokenizer
s=["[CLS]","[PAD]","[SEP]","[UNK]","[MASK]"]
wpt=BertWordPieceTokenizer(strip_accents=False)
wpt.train(files=["train.txt"],vocab_size=8000,special_tokens=s)
wpt.save_model(".")
tkz=BertTokenizerFast(vocab_file="vocab.txt",strip_accents=False,never_split=s)
tkz.save_pretrained("my-dir/tokenizer-nl")
```

図 64: オランダ語向けトークナイザの作成

次に train.txt の文章から、DeBERTa モデルを my-dir/deberta-nl に作成しましょう。図 65 では、入出力幅 128 トークン・入出力ベクトル 256 次元・深さ 12 層・アテンションヘッド 4 個・中間ベクトル 768 次元という、やや小さめのモデルにしており、CPU なら 1 時間程度、GPU なら 10 分程度で作成できます。

最後に、my-dir/deberta-nl と UD_Dutch-Alpino・UD_Dutch-LassySmall から、品詞付与・係り受け解析モデル my-dir/deberta-nl-upos を製作しましょう (図 66)。うまく my-dir/deberta-nl-upos が製作できたら、esupar でテストしてみましょう (図 67)。

^[83]https://github.com/UniversalDependencies/UD_Dutch-Alpino

^[84]https://github.com/UniversalDependencies/UD_Dutch-LassySmall

```

from transformers import (AutoTokenizer,DebertaV2Config,DebertaV2ForMaskedLM,
    DataCollatorForLanguageModeling,TrainingArguments,Trainer)
tkz=AutoTokenizer.from_pretrained("my-dir/tokenizer-nl",model_max_length=128)
t=tkz.convert_tokens_to_ids(["[CLS]","[PAD]","[SEP]","[UNK]","[MASK]"])
cfg=DebertaV2Config(hidden_size=256,num_hidden_layers=12,num_attention_heads=4,
    intermediate_size=768,relative_attention=True,position_biased_input=False,
    pos_att_type=["p2c","c2p"],max_position_embeddings=tkz.model_max_length,
    vocab_size=len(tkz),tokenizer_class=type(tkz).__name__,
    bos_token_id=t[0],pad_token_id=t[1],eos_token_id=t[2])
arg=TrainingArguments(report_to="none",output_dir="/tmp",save_total_limit=2,
    num_train_epochs=3,per_device_train_batch_size=64)
class ReadLineDS(object):
    def __init__(self,file,tokenizer):
        self.tokenizer=tokenizer
        with open(file,"r",encoding="utf-8") as r:
            self.lines=[s.strip() for s in r if s.strip()!=""]
    __len__=lambda self:len(self.lines)
    __getitem__=lambda self,i:self.tokenizer(self.lines[i],truncation=True,
        add_special_tokens=True,max_length=self.tokenizer.model_max_length-2)
trn=Trainer(args=arg,data_collator=DataCollatorForLanguageModeling(tkz),
    model=DebertaV2ForMaskedLM(cfg),train_dataset=ReadLineDS("train.txt",tkz))
trn.train()
trn.save_model("my-dir/deberta-nl")
tkz.save_pretrained("my-dir/deberta-nl")

```

図 65: オランダ語ミニ DeBERTa モデルの作成

```
!python -m esupar.train my-dir/deberta-nl my-dir/deberta-nl-upos .
```

図 66: オランダ語向け品詞付与・係り受け解析ミニモデルの製作

```

!pip install esupar
import esupar
nlp=esupar.load("my-dir/deberta-nl-upos")
doc=nlp("In de ziekenkamer zag hij ziek uit maar wilde goed uitzien")
import deplacy
deplacy.serve(doc,port=None)

```

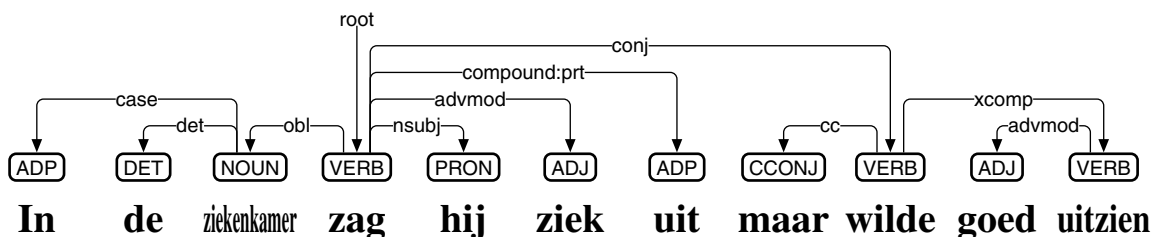


図 67: オランダ語向け品詞付与・係り受け解析ミニモデルのテスト

8 タイ語編

8.1 タイ語 UD における品詞付与

書写言語としてのタイ語は、単語と単語の間に区切りがありません。数字や氏名の前後には、空白を入れる場合があるのですが、単語の区切りというわけではないのです。事前学習モデルでタイ語を扱う際には、各トークンを単語と合致させるのが理想です。しかし、それは現実には困難なため、各トークンは単語より短めにトークナイズした上で、系列ラベリング (UPOS 付与) によって、単語をまとめ上げる手法が考えられます。すなわち、トークンが単語と合致した場合には UPOS をラベルとして付与し、そうでない場合には、最初のトークンに「B-」を付けた UPOS を、それ以降のトークンに「I-」を付けた UPOS を、それぞれ付与するのです。

```
!pip install deplacy transformers
class SeqL(object):
    def __init__(self,bert):
        from transformers import pipeline
        self.tagger=pipeline(task="ner",model=bert)
    def __call__(self,text):
        w=[(t["start"],t["end"],t["entity"],t["word"]) for t in self.tagger(text)]
        u,z="# text = "+text.replace("\n"," ")+"\n","_"
        for i,(s,e,p,t) in enumerate(w,1):
            m,q=z if i<len(w) and e<w[i][0] else "SpaceAfter=No",p.split("|")
            f=z if p.find("=")<0 else "|".join(t for t in q if t.find("=")>0)
            u+="\t".join([str(i),t.strip(),z,q[0].replace(".",","),z,f,z,z,z,m])+"\n"
        return u+"\n"
nlp=SeqL("KoichiYasuoka/modernbert-base-thai-cc100-ud-embeds")
doc=nlp("หลายหัวดีกว่าหัวเดียว")
import deplacy
deplacy.serve(doc,port=None)
```

DET	NOUN	VERB	ADP	NOUN	B-NUM	I-NUM
หลาย	หัว	ดี	กว่า	หัว	เ	เดียว

図 68: modernbert-base-thai-cc100-ud-embeds による UPOS 付与

Google Colaboratory 上の Transformers^[26]を用いて、タイ語の例文「หลายหัวดีกว่าหัวเดียว」に modernbert-base-thai-cc100-ud-embeds^[85]で UPOS を付与するプログラムと、その結果を図 68 に示します。「หลาย」「หัว」「ดี」「กว่า」の各単語では、各トークンがそのまま単語と合致しており、それぞれ UPOS 品詞 DET・NOUN・VERB・ADP が付与されています。「เดียว」は、2つのトークンに泣き別れてしまっており、それぞれに B-NUM と I-NUM が付与されています。なお「ดี」は形容詞とみなせますが、タイ語においては形容詞と動詞に文法的な区別がないことから、ADJ ではなく VERB が付与されています。

^[85]<https://huggingface.co/KoichiYasuoka/modernbert-base-thai-cc100-ud-embeds>

8.2 タイ語 UD における係り受け解析

modernbert-base-thai-cc100-ud-embeds は、隣接行列型係り受け解析アルゴリズムを、事前学習モデルで実装しています。Google Colaboratory 上の Transformers を用いて、タイ語の例文「หลายหัวดีกว่าหัวเดียว」を modernbert-base-thai-cc100-ud-embeds で係り受け解析するプログラムと、その結果を図 69 に示します。解析結果を見ていくことにしましょう。

```
!pip install deplacy transformers
from transformers import pipeline
mdl="KoichiYasuoka/modernbert-base-thai-cc100-ud-embeds"
nlp=pipeline("universal-dependencies",mdl,trust_remote_code=True)
doc=nlp("หลายหัวดีกว่าหัวเดียว")
import deplacy
deplacy.serve(doc,port=None)
```

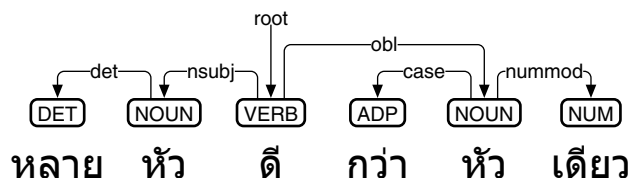


図 69: modernbert-base-thai-cc100-ud-embeds によるタイ語の係り受け解析

動詞(形容詞)「ดี」には、名詞「หัว」が2つ繋がっており、前者が主語(nsubj)、後者が斜格補語(obl)です。2つの「หัว」は、それぞれ数量詞「หลาย」「เดียว」で修飾されており、前者は名詞に前置されていて det で、後者は名詞に後置されていて nummod で繋がっています。後者の「หัว」には、前置詞「กว่า」が case で繋がられています。

8.3 タイ語ミニ DeBERTa モデルの製作

Google Colaboratory 上の esupar で、タイ語ミニ DeBERTa モデルを製作してみましょう。手順を図 70～73 に示します。

最初に、訓練のための文章を準備しましょう(図 70)。train.txt の文章の材料として、ここでは VISTEC-TP-TH-2021^[86]と UD_Thai-TUD^[87]を元にしていますが、好みに応じ、タイ語の他の文章を増量するのもいいでしょう。train.txt と同時に、各単語を空白で区切った token.txt も準備しており、トークナイザの作成に用います。また、UD_Thai-TUD の各 CoNLL-U ファイルは、品詞付与・係り受け解析の学習にも用います。

続けて、token.txt の各単語から、タイ語向けトークナイザを、my-dir/tokenizer-th に作成しましょう。図 71 では、Unigram トークナイザ(SentencePiece^[88]の改造版)を token.txt

^[86]Peerat Limkonchotiwat, Wannaphong Phatthiyaphaibun, Raheem Sarwar, Ekapol Chuangsuwanich, Sarana Nutanong: Handling Cross- and Out-of-Domain Samples in Thai Word Segmentaion, Findings of the Association for Computational Linguistics: ACL-IJCNLP 2021 (August 2021), pp.1003-1016.

^[87]Panyut Sriwirote, Wei Qi Leong, Charin Polpanumas, Santhawat Thanyawong, William Chandra Tjhi, Wirote Aroonmanakun, Attapol T. Rutherford: The Thai Universal Dependency Treebank, Transactions of the Association for Computational Linguistics, Vol.13, pp.376-391 (July 2025).

^[88]Taku Kudo, John Richardson: SentencePiece: A simple and language independent subword tokenizer and detokenizer for Neural Text Processing, Proceedings of the 2018 Conference on Empirical Methods in Natural Language Processing (November 2018): System Demonstrations, pp.66-71.

```

!pip install esupar spacy-alignments accelerate
import os
url="https://github.com/mrpeerat/OSkut"
v=os.path.join(os.path.basename(url),"VISTEC-TP-TH-2021")
!test -d {v} || git clone --depth=1 {url}
!sed "s/<[^>]*>//g;s/[|_]/ /g" {v}/*/*processed.txt > token.txt
url="https://github.com/UniversalDependencies/UD_Thai-TUD"
d=os.path.basename(url)
!test -d {d} || git clone --depth=1 {url}
!for F in train dev test ; do cp {d}/*-$$F.conllu $$F.conllu ; done
s='' '{if(NF==10&&$1~/^[1-9][0-9]*$/)}printf($1>1?" %s":"%s",$2)>>"token.txt";
  if(NF==0)print>>"token.txt";if($0~/^# text = /)print substr($0,10)}''
!nawk -F'\t' '{s}' *.conllu | cat {v}/*/*raw.txt - > train.txt

```

図 70: タイ語ミニモデル製作のための準備

に適用し、それを `DebertaV2TokenizerFast` という Transformers のトークナイザに組み上げて、タイ語向けトークナイザを作成しています。語彙数を 3000 に絞った上で、各トークンの文字数を最大 4 文字 (母音や声調記号を 1 文字に数える) としています。なお、各トークンは、タイ文字クラスター (คลัสเตอร์อักษรไทย)^{[89][90]} を途中で切ってしまうような調整しています。

```

from transformers import DebertaV2TokenizerFast
from tokenizers import (Tokenizer,models,pre_tokenizers,normalizers,processors,
    decoders,trainers)
import json,unicodedata
s=["[CLS]","[PAD]","[SEP]","[UNK]","[MASK]"]
spt=Tokenizer(models.Unigram())
spt.pre_tokenizer=pre_tokenizers.Sequence([pre_tokenizers.Whitespace(),
    pre_tokenizers.Punctuation()])
spt.normalizer=normalizers.Sequence([normalizers.Nmt(),normalizers.NFKC()])
spt.post_processor=processors.TemplateProcessing(single="[CLS] $A [SEP]",
    pair="[CLS] $A [SEP] $B:1 [SEP]:1",special_tokens=[("[CLS]",0),("[SEP]",2)])
spt.decoder=decoders.WordPiece(prefix="",cleanup=True)
spt.train(trainer=trainers.UnigramTrainer(vocab_size=3000,max_piece_length=4,
    special_tokens=s,unk_token="[UNK]",n_sub_iterations=2),files=["token.txt"])
d=json.loads(spt.to_str())
d["model"]["vocab"]=[t for t in d["model"]["vocab"] if len(t[0])<2 or
    unicodedata.category(t[0][0])!="Mn" and int((ord(t[0][-1])-1)/7)!=521]
spt.from_str(json.dumps(d)).save("tokenizer.json")
tkz=DebertaV2TokenizerFast(tokenizer_file="tokenizer.json",split_by_punct=True,
    do_lower_case=False,keep_accents=True,vocab_file="/dev/null")
tkz.save_pretrained("my-dir/tokenizer-th")

```

図 71: タイ語向けトークナイザの作成

^[89]วิรัช ศรีเลิศล้ำวานิช: การตัดคำไทยในระบบแปลภาษา [Word Segmentation for Thai in Machine Translation System], การแปลภาษาด้วยคอมพิวเตอร์ [Machine Translation], Bangkok: NECTEC (1993), pp.556-561.

^[90]Thanaruk Theeramunkong, Virach Sornlertlamvanich, Thanasan Tanhermhong, Wirat Chinnan: Character Cluster Based Thai Information Retrieval, IRAL '00: Proceedings of the Fifth International Workshop on Information Retrieval with Asian Languages (November 2000), pp.75-80

次に train.txt の文章から、DeBERTa モデルを my-dir/deberta-th に作成しましょう。図 72 では、入出力幅 128 トークン・入出力ベクトル 256 次元・深さ 12 層・アテンションヘッド 4 個・中間ベクトル 768 次元という、やや小さめのモデルにしていますが、GPU でも 1 時間程度かかってしまいます。

最後に、my-dir/deberta-th と UD_Thai-TUD から、品詞付与・係り受け解析モデル my-dir/deberta-th-upos を製作しましょう (図 73)。うまく my-dir/deberta-th-upos ができたら、esupar でテストしてみましょう (図 74)。

```

from transformers import (AutoTokenizer,DebertaV2Config,DebertaV2ForMaskedLM,
    DataCollatorForLanguageModeling,TrainingArguments,Trainer)
tkz=AutoTokenizer.from_pretrained("my-dir/tokenizer-th",model_max_length=128)
t=tkz.convert_tokens_to_ids(["[CLS]","[PAD]","[SEP]","[UNK]","[MASK]"])
cfg=DebertaV2Config(hidden_size=256,num_hidden_layers=12,num_attention_heads=4,
    intermediate_size=768,relative_attention=True,position_biased_input=False,
    pos_att_type=["p2c","c2p"],max_position_embeddings=tkz.model_max_length,
    vocab_size=len(tkz),tokenizer_class=type(tkz).__name__,
    bos_token_id=t[0],pad_token_id=t[1],eos_token_id=t[2])
arg=TrainingArguments(report_to="none",output_dir="/tmp",save_total_limit=2,
    num_train_epochs=3,per_device_train_batch_size=64)
class ReadLineDS(object):
    def __init__(self,file,tokenizer):
        self.tokenizer=tokenizer
        with open(file,"r",encoding="utf-8") as r:
            self.lines=[s.strip() for s in r if s.strip()!=""]
    __len__=lambda self:len(self.lines)
    __getitem__=lambda self,i:self.tokenizer(self.lines[i],truncation=True,
        add_special_tokens=True,max_length=self.tokenizer.model_max_length-2)
trn=Trainer(args=arg,data_collator=DataCollatorForLanguageModeling(tkz),
    model=DebertaV2ForMaskedLM(cfg),train_dataset=ReadLineDS("train.txt",tkz))
trn.train()
trn.save_model("my-dir/deberta-th")
tkz.save_pretrained("my-dir/deberta-th")

```

図 72: タイ語ミニ DeBERTa モデルの作成

```

s,d="my-dir/deberta-th","my-dir/deberta-th-upos"
!python -m esupar.train {s} {d} 32 /tmp train.upos
!python -m esupar.train {d} {d} 32 /// train.conllu dev.conllu test.conllu

```

図 73: タイ語向け品詞付与・係り受け解析ミニモデルの製作

```

!pip install esupar
import esupar
nlp=esupar.load("my-dir/deberta-th-upos")
doc=nlp("หลายหัวดีกว่าหัวเดียว")
import deplacy
deplacy.serve(doc,port=None)

```

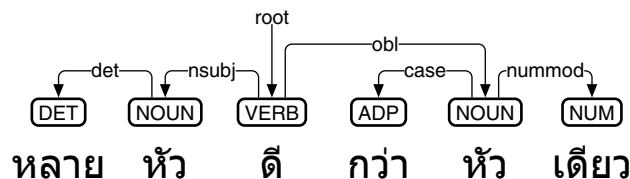


図 74: タイ語向け品詞付与・係り受け解析ミニモデルのテスト

9 ベトナム語編

9.1 ベトナム語UDにおける品詞付与

書写言語としてのベトナム語は、単語と単語の間ではなく、音節と音節の間に空白があります。したがって、複数の音節で構成される単語は、語中に空白を含むことになります。事前学習モデルでベトナム語を扱う際には、各トークンを音節と合致させて、系列ラベリング (UPOS 付与) によって、単語をまとめ上げる手法が考えられます。すなわち、単音節の単語に対しては UPOS をラベルとして付与し、そうでない場合には、最初のトークンに「B-」を付けた UPOS を、それ以降のトークンに「I-」を付けた UPOS を、それぞれ付与するのです。

```
!pip install deplacy transformers
class SeqL(object):
    def __init__(self,bert):
        from transformers import pipeline
        self.tagger=pipeline(task="ner",model=bert)
    def __call__(self,text):
        w=[(t["start"],t["end"],t["entity"],t["word"]) for t in self.tagger(text)]
        u,z="# text = "+text.replace("\n"," ")+"\n", "_"
        for i,(s,e,p,t) in enumerate(w,1):
            m,q=z if i<len(w) and e<w[i][0] else "SpaceAfter=No",p.split("|")
            f=z if p.find("=")<0 else "|".join(t for t in q if t.find("=")>0)
            u+="%t".join([str(i),t.strip(),z,q[0].replace(".",""),z,f,z,z,z,m])+"\n"
        return u+"\n"
nlp=SeqL("KoichiYasuoka/roberta-base-vietnamese-upos")
doc=nlp("Nó không xử lý được các từ có nhiều âm tiết")
import deplacy
deplacy.serve(doc,port=None)
```

PRON ADV B-VERB I-VERB ADV DET NOUN VERB ADJ B-NOUN I-NOUN
Nó không xử lý được các từ có nhiều âm tiết

図 75: roberta-base-vietnamese-upos による UPOS 付与

Google Colaboratory 上の Transformers^[26]を用いて、roberta-base-vietnamese-upos^[91]でベトナム語の例文「Nó không xử lý được các từ có nhiều âm tiết」に UPOS を付与するプログラムと、その結果を図 75 に示します。「Nó」「không」「được」「các」「từ」「có」「nhiều」の各単語では、各トークンがそのまま単語と合致しており、それぞれ UPOS 品詞 PRON・ADV・ADV・DET・NOUN・VERB・ADJ が付与されています。「xử lý」と「âm tiết」は 2 音節の単語であり、「xử lý」には B-VERB と I-VERB が、「âm tiết」には B-NOUN と I-NOUN が、それぞれ付与されています。

^[91]<https://huggingface.co/KoichiYasuoka/roberta-base-vietnamese-upos>

9.2 ベトナム語 UD における係り受け解析

esupar^[47]は、隣接行列型係り受け解析アルゴリズムを、事前学習モデルで実装しています。Google Colaboratory 上の esupar を用いて、roberta-base-vietnamese-upos でベトナム語の例文「Nó không xử lý được các từ có nhiều âm tiết」を係り受け解析するプログラムと、その結果を図 76 に示します。解析結果を見ていくことにしましょう。

```
!pip install esupar
import esupar
nlp=esupar.load("KoichiYasuoka/roberta-base-vietnamese-upos")
doc=nlp("Nó không xử lý được các từ có nhiều âm tiết")
import deplacy
deplacy.serve(doc,port=None)
```

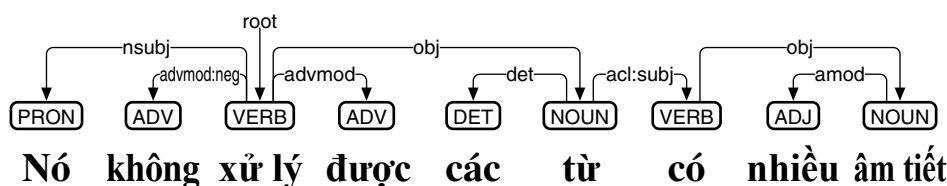


図 76: roberta-base-vietnamese-upos によるベトナム語の係り受け解析

動詞「xử lý」の主語「Nó」は **nsubj** で、目的語「các từ」は **obj** で繋がられています。ただし「các từ」は 2 語だとみなされており、内部に **det** が繋がられています。「không」「được」は「xử lý」を否定しており、それぞれ **advmod:neg** と **advmod** で繋がられています。「có nhiều âm tiết」は「các từ」の連体修飾節だとみなされていて、**acl:subj** で繋がられています。「có nhiều âm tiết」の内部では、「có」の目的語「âm tiết」が **obj** で繋がっていて、さらに「nhiều」が **amod** で繋がられています。

なお、ベトナム語 UD の係り受けタグには、表 4 と **advmod:neg**・**acl:subj** に加え、数多くの拡張が用いられている^[92]ののですが、拡張どうしの中で混乱が起こっており、事態の收拾が望まれます。

9.3 ベトナム語ミニ DeBERTa モデルの製作

Google Colaboratory 上の esupar で、ベトナム語ミニ DeBERTa モデルを製作してみましょう。手順を図 77～80 に示します。

最初に、訓練のための文章を準備しましょう (図 77)。train.txt の文章の材料として、ここでは VNESEcorpus^[93]と UD_Vietnamese-VTB^[94]を元にしていますが、好みに応じ、ベトナム語の他の文章を増量するのもいいでしょう。

続けて、ベトナム語向けトークナイザを、my-dir/tokenizer-vi に作成しましょう。図 78 では、Transformers の BertTokenizerFast を使って、ベトナム語向けトークナイザを作成しています。

^[92]https://universaldependencies.org/treebanks/vi_vtb/#relations-overview

^[93]Tuan Anh Luu, Kazuhide Yamamoto: A Pointwise Approach for Vietnamese Diacritics Restoration, Proceedings of the 2012 International Conference on Asian Language Processing (November 2012), pp.189-192.

^[94]https://github.com/UniversalDependencies/UD_Vietnamese-VTB

```

!pip install esupar accelerate
import os
url="https://github.com/UniversalDependencies/UD_Vietnamese-VTB"
d=os.path.basename(url)
!test -d {d} || git clone --depth=1 {url}
!for F in train dev test ; do cp {d}/*-$$F.conllu $$F.conllu ; done
url="https://filedn.com/lit4DCIlHwxfS1gj9zcYuDJ/SNOW/VNESEcorpus.txt"
f=os.path.basename(url)
!test -f {f} || curl -LO {url}
!sed -n "s/^# text = */p" *.conllu | cat {f} - > train.txt

```

図 77: ベトナム語ミニモデル製作のための準備

```

from transformers import BertTokenizerFast
from tokenizers import BertWordPieceTokenizer
s=["[CLS]","[PAD]","[SEP]","[UNK]","[MASK]"]
wpt=BertWordPieceTokenizer(strip_accents=False)
wpt.train(files=["train.txt"],vocab_size=8000,special_tokens=s)
wpt.save_model(".")
tkz=BertTokenizerFast(vocab_file="vocab.txt",strip_accents=False,never_split=s)
tkz.save_pretrained("my-dir/tokenizer-vi")

```

図 78: ベトナム語向けトークナイザの作成

```

from transformers import (AutoTokenizer,DebertaV2Config,DebertaV2ForMaskedLM,
    DataCollatorForLanguageModeling,TrainingArguments,Trainer)
tkz=AutoTokenizer.from_pretrained("my-dir/tokenizer-vi",model_max_length=128)
t=tkz.convert_tokens_to_ids(["[CLS]","[PAD]","[SEP]","[UNK]","[MASK]"])
cfg=DebertaV2Config(hidden_size=256,num_hidden_layers=12,num_attention_heads=4,
    intermediate_size=768,relative_attention=True,position_biased_input=False,
    pos_att_type=["p2c","c2p"],max_position_embeddings=tkz.model_max_length,
    vocab_size=len(tkz),tokenizer_class=type(tkz).__name__,
    bos_token_id=t[0],pad_token_id=t[1],eos_token_id=t[2])
arg=TrainingArguments(report_to="none",output_dir="/tmp",save_total_limit=2,
    num_train_epochs=3,per_device_train_batch_size=64)
class ReadLineDS(object):
    def __init__(self,file,tokenizer):
        self.tokenizer=tokenizer
        with open(file,"r",encoding="utf-8") as r:
            self.lines=[s.strip() for s in r if s.strip()!=""]
    __len__=lambda self:len(self.lines)
    __getitem__=lambda self,i:self.tokenizer(self.lines[i],truncation=True,
        add_special_tokens=True,max_length=self.tokenizer.model_max_length-2)
trn=Trainer(args=arg,data_collator=DataCollatorForLanguageModeling(tkz),
    model=DebertaV2ForMaskedLM(cfg),train_dataset=ReadLineDS("train.txt",tkz))
trn.train()
trn.save_model("my-dir/deberta-vi")
tkz.save_pretrained("my-dir/deberta-vi")

```

図 79: ベトナム語ミニ DeBERTa モデルの作成

次に train.txt の文章から、DeBERTa モデルを my-dir/deberta-vi に作成しましょう。図 79 では、入出力幅 128 トークン・入出力ベクトル 256 次元・深さ 12 層・アテンションヘッド 4 個・中間ベクトル 768 次元という、やや小さめのモデルにしていますが、GPU でも 1 時間程度かかってしまいます。

最後に、my-dir/deberta-vi と UD_Vietnamese-VTB から、品詞付与・係り受け解析モデル my-dir/deberta-vi-upos を製作しましょう (図 80)。うまく my-dir/deberta-vi-upos ができたら、esupar でテストしてみましょう (図 81)。

```
!python -m esupar.train my-dir/deberta-vi my-dir/deberta-vi-upos .
```

図 80: ベトナム語向け品詞付与・係り受け解析ミニモデルの製作

```
!pip install esupar
import esupar
nlp=esupar.load("my-dir/deberta-vi-upos")
doc=nlp("Nó không xử lý được các từ có nhiều âm tiết")
import deplacy
deplacy.serve(doc,port=None)
```

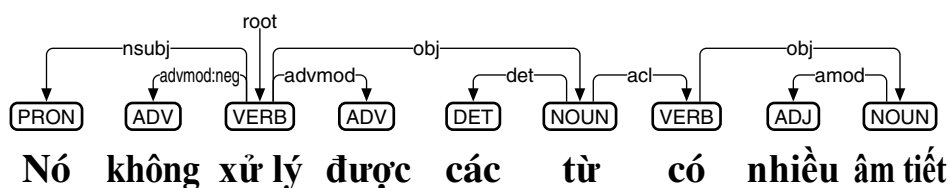


図 81: ベトナム語向け品詞付与・係り受け解析ミニモデルのテスト

10 ロシア語編

10.1 ロシア語 UD における品詞付与と係り受け解析

modernbert-base-russian-ud-embeds^[95]は、隣接行列型係り受け解析アルゴリズムを、事前学習モデルで実装しています。Google Colaboratory 上の Transformers^[26]を用いて、modernbert-base-russian-ud-embeds でロシア語の例文「Москва слезам не верила а верила любви」を品詞付与・係り受け解析するプログラムと、その結果を図 82 に示します。解析結果を見ていくことにしましょう。

```
!pip install deplacy transformers
from transformers import pipeline
mdl="KoichiYasuoka/modernbert-base-russian-ud-embeds"
nlp=pipeline("universal-dependencies",mdl,trust_remote_code=True)
doc=nlp("Москва слезам не верила а верила любви ")
import deplacy
deplacy.serve(doc,port=None)
```

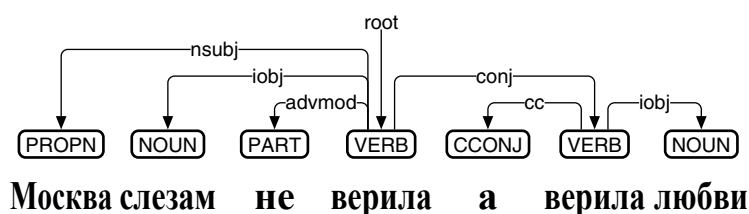


図 82: modernbert-base-russian-ud-embeds によるロシア語の係り受け解析

2つの動詞「верила」が並置されていて、conjで繋がられています。前者には、否定の「не」がadvmodで繋がられています。後者には、接続詞「а」がccで繋がられています。「Москва」は、両方の「верила」の主語ですが、前者にnsubjで繋がられています。前者の目的語は「слезам」、後者の目的語は「любви」で、いずれもobjではなくioobjで繋がられています。Trankitは、与格(дательный)の目的語を、間接目的語(ioobj)だとみなしてしまうようです。

なお、ロシア語 UD の係り受けタグには、表 4 に加え、受動態を表す nsubj:pass・csubj:pass・aux:pass、受動態の動作主を表す obl:agent、関係代名詞による連体修飾節を表す acl:relcl、数詞と名詞の格変化の関係を示す nummod:gov・nummod:entity、氏名を表す flat:name、ロシア語以外を表す flat:foreign も用いられます^[96]。

10.2 ロシア語ミニ DeBERTa モデルの製作

Google Colaboratory 上の esupar^[47]で、ロシア語ミニ DeBERTa モデルを製作してみましょう。手順を図 83～86 に示します。

^[95]<https://huggingface.co/KoichiYasuoka/modernbert-base-russian-ud-embeds>

^[96]<https://universaldependencies.org/ru/#relations-overview>

```

!pip install esupar spacy-alignments accelerate
import os
url="https://github.com/UniversalDependencies/UD_Russian-SynTagRus"
!test -d {os.path.basename(url)} || git clone --depth=1 {url}
url="https://github.com/UniversalDependencies/UD_Russian-Taiga"
!test -d {os.path.basename(url)} || git clone --depth=1 {url}
!for F in train dev test ; do cat UD_Russian-*/*-$$F*.conllu > $$F.conllu ; done
!sed -n "s/^# text = //p" *.conllu > train.txt

```

図 83: ロシア語ミニモデル製作のための準備

```

from transformers import BertTokenizerFast
from tokenizers import BertWordPieceTokenizer
s=["[CLS]","[PAD]","[SEP]","[UNK]","[MASK]"]
wpt=BertWordPieceTokenizer(strip_accents=False)
wpt.train(files=["train.txt"],vocab_size=30000,special_tokens=s)
wpt.save_model(".")
tkz=BertTokenizerFast(vocab_file="vocab.txt",strip_accents=False,never_split=s)
tkz.save_pretrained("my-dir/tokenizer-ru")

```

図 84: ロシア語向けトークナイザの作成

```

from transformers import (AutoTokenizer,DebertaV2Config,DebertaV2ForMaskedLM,
    DataCollatorForLanguageModeling,TrainingArguments,Trainer)
tkz=AutoTokenizer.from_pretrained("my-dir/tokenizer-ru",model_max_length=128)
t=tkz.convert_tokens_to_ids(["[CLS]","[PAD]","[SEP]","[UNK]","[MASK]"])
cfg=DebertaV2Config(hidden_size=256,num_hidden_layers=12,num_attention_heads=4,
    intermediate_size=768,relative_attention=True,position_biased_input=False,
    pos_att_type=["p2c","c2p"],max_position_embeddings=tkz.model_max_length,
    vocab_size=len(tkz),tokenizer_class=type(tkz).__name__,
    bos_token_id=t[0],pad_token_id=t[1],eos_token_id=t[2])
arg=TrainingArguments(report_to="none",output_dir="/tmp",save_total_limit=2,
    num_train_epochs=3,per_device_train_batch_size=64)
class ReadLineDS(object):
    def __init__(self,file,tokenizer):
        self.tokenizer=tokenizer
        with open(file,"r",encoding="utf-8") as r:
            self.lines=[s.strip() for s in r if s.strip()!=""]
    __len__=lambda self:len(self.lines)
    __getitem__=lambda self,i:self.tokenizer(self.lines[i],truncation=True,
        add_special_tokens=True,max_length=self.tokenizer.model_max_length-2)
trn=Trainer(args=arg,data_collator=DataCollatorForLanguageModeling(tkz),
    model=DebertaV2ForMaskedLM(cfg),train_dataset=ReadLineDS("train.txt",tkz))
trn.train()
trn.save_model("my-dir/deberta-ru")
tkz.save_pretrained("my-dir/deberta-ru")

```

図 85: ロシア語ミニ DeBERTa モデルの作成

最初に、訓練のための文章を準備しましょう (図 83)。train.txt の文章の材料として、ここでは UD_Russian-SynTagRus^[97] と UD_Russian-Taiga^[98] を元にしてありますが、好みに応じ、ロシア語の他の文章を増量するのもいいでしょう。なお、各 CoNLL-U ファイルは、品詞付与・係り受け解析の学習にも用います。

続けて、ロシア語向けトークナイザを、my-dir/tokenizer-ru に作成しましょう。図 84 では、Transformers^[26] の BertTokenizerFast を使って、ロシア語向けトークナイザを作成しています。

次に train.txt の文章から、DeBERTa モデルを my-dir/deberta-ru に作成しましょう。図 85 では、入出力幅 128 トークン・入出力ベクトル 256 次元・深さ 12 層・アテンションヘッド 4 個・中間ベクトル 768 次元という、やや小さめのモデルにしています。

最後に、my-dir/deberta-ru と、UD_Russian-SynTagRus・UD_Russian-Taiga の CoNLL-U ファイルから、品詞付与・係り受け解析モデル my-dir/deberta-ru-upos を製作しましょう (図 86)。うまく my-dir/deberta-ru-upos が製作できたら、esupar でテストしてみましょう (図 87)。

```
!python -m esupar.train my-dir/deberta-ru my-dir/deberta-ru-upos .
```

図 86: ロシア語向け品詞付与・係り受け解析ミニモデルの製作

```
!pip install esupar
import esupar
nlp=esupar.load("my-dir/deberta-ru-upos")
doc=nlp("Москва слезам не верила а верила любви")
import deplacy
deplacy.serve(doc,port=None)
```

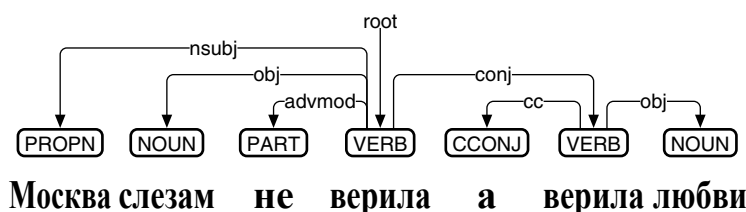


図 87: ロシア語向け品詞付与・係り受け解析ミニモデルのテスト

^[97]Kira Drogonova, Olga Lyashevskaya, Daniel Zeman: Data Conversion and Consistency of Monolingual Corpora: Russian UD Treebanks, Proceedings of the 17th International Workshop on Treebanks and Linguistic Theories (December 2018), pp.52-65.

^[98]T. O. Sharvina, O. Shapovalova: To the Methodology of Corpus Construction for Machine Learning: “Taiga” Syntax Tree Corpus and Parser, Proceedings of «CORPUS LINGUISTICS-2017» (June 2017), pp.78-84.

11 ウクライナ語編

11.1 ウクライナ語 UD における品詞付与と係り受け解析

Trankit^[64]は、XLM-RoBERTa をベースに、隣接行列型係り受け解析アルゴリズムを、様々な言語向けに実装しています。Google Colaboratory 上の Trankit を用いて、ウクライナ語の例文「Не скупись на втіху їй і ласку любий брате」を品詞付与・係り受け解析するプログラムと、その結果を図 88 に示します。解析結果を見ていくことにしましょう。

```
!pip install git+https://github.com/KoichiYasuoka/trankit deplacy
import trankit
nlp=trankit.Pipeline("ukrainian")
doc=nlp("Не скупись на втіху їй і ласку любий брате")
import deplacy
deplacy.serve(doc,port=None)
```

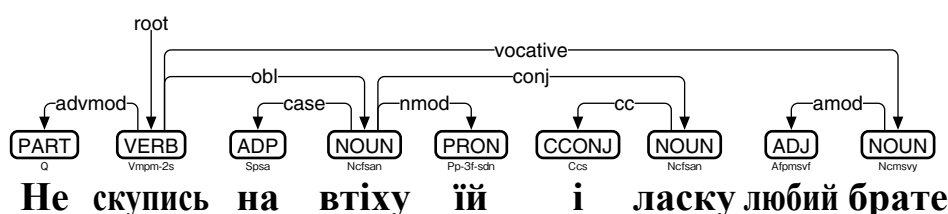


図 88: Trankit によるウクライナ語の係り受け解析

不完了相動詞「скупись」は、否定の「Не」が **advmod** で繋がれていて、禁止の命令を表しています。名詞「брате」が **vocative** で繋がれていて、さらに「любий」が **amod** で繋がれています。「втіху」と「ласку」は、いずれも「скупись」の斜格補語で、たがいに **conj** で繋がれていると同時に、前者が「скупись」に **obl** で繋がれています。前置詞「на」と代名詞「їй」は、前者にそれぞれ **case** と **nmod** で繋がれており、接続詞「і」は後者に **cc** で繋がれています。

なお、ウクライナ語 UD の係り受けタグには、表 4 に加え、数多くの拡張が用いられている^[99]のようですが、拡張どうしの中で混乱が起こっており、事態の收拾が望まれます。

11.2 ウクライナ語ミニ DeBERTa モデルの製作

Google Colaboratory 上の esupar^[47]で、ウクライナ語ミニ DeBERTa モデルを製作してみましょう。手順を図 89～92 に示します。

最初に、訓練のための文章を準備しましょう (図 89)。train.txt の文章の材料として、ここでは UD_Ukrainian-IU^[100]と Ukr-Synth^[101]を元にしてしていますが、好みに応じ、ウクライナ語の他の文章を増量するのもいいでしょう。なお、UD_Ukrainian-IU・Ukr-Synth の各 CoNLL-U ファイルは、train.upos にまとめ上げて、品詞付与の学習に用います。また、UD_Ukrainian-IU の各 CoNLL-U ファイルは、係り受け解析の学習にも用います。

^[99]https://universaldependencies.org/treebanks/uk_iu/#relations-overview

^[100]https://github.com/UniversalDependencies/UD_Ukrainian-IU

^[101]<https://huggingface.co/datasets/ukr-models/Ukr-Synth>

```

!pip install esupar spacy-alignments accelerate
import os
url="https://github.com/UniversalDependencies/UD_Ukrainian-IU"
d=os.path.basename(url)
!test -d {d} || git clone --depth=1 {url}
!for F in train dev test ; do cp {d}/*-$$F.conllu $$F.conllu ; done
url="https://huggingface.co/datasets/ukr-models/Ukr-Synth"
d=os.path.basename(url)
!test -d {d} || git-lfs clone --depth=1 {url}
!zcat {d}/data/*.conllu.gz | tr -d "\015" | cat - *.conllu > train.upos
!sed -n "s/# text = //p" train.upos > train.txt

```

図 89: ウクライナ語ミニモデル製作のための準備

```

from transformers import BertTokenizerFast
from tokenizers import BertWordPieceTokenizer
s=["[CLS]","[PAD]","[SEP]","[UNK]","[MASK]"]
wpt=BertWordPieceTokenizer(strip_accents=False)
wpt.train(files=["train.txt"],vocab_size=30000,special_tokens=s)
wpt.save_model(".")
tkz=BertTokenizerFast(vocab_file="vocab.txt",strip_accents=False,never_split=s)
tkz.save_pretrained("my-dir/tokenizer-uk")

```

図 90: ウクライナ語向けトークナイザの作成

続けて、ウクライナ語向けトークナイザを、my-dir/tokenizer-uk に作成しましょう。図 90 では、Transformers^[26]の BertTokenizerFast を使って、ウクライナ語向けトークナイザを作成しています。

次に train.txt の文章から、DeBERTa モデルを my-dir/deberta-uk に作成しましょう。図 91 では、入出力幅 128 トークン・入出力ベクトル 256 次元・深さ 12 層・アテンションヘッド 4 個・中間ベクトル 768 次元という、やや小さめのモデルにしています。

最後に、my-dir/deberta-uk と UD_Ukrainian-IU・Ukr-Synth から、品詞付与・係り受け解析モデル my-dir/deberta-uk-upos を製作しましょう (図 92)。うまく my-dir/deberta-uk-upos が製作できたら、esupar でテストしてみましょう (図 93)。

```

from transformers import (AutoTokenizer,DebertaV2Config,DebertaV2ForMaskedLM,
    DataCollatorForLanguageModeling,TrainingArguments,Trainer)
tkz=AutoTokenizer.from_pretrained("my-dir/tokenizer-uk",model_max_length=128)
t=tkz.convert_tokens_to_ids(["[CLS]","[PAD]","[SEP]","[UNK]","[MASK]"])
cfg=DebertaV2Config(hidden_size=256,num_hidden_layers=12,num_attention_heads=4,
    intermediate_size=768,relative_attention=True,position_biased_input=False,
    pos_att_type=["p2c","c2p"],max_position_embeddings=tkz.model_max_length,
    vocab_size=len(tkz),tokenizer_class=type(tkz).__name__,
    bos_token_id=t[0],pad_token_id=t[1],eos_token_id=t[2])
arg=TrainingArguments(report_to="none",output_dir="/tmp",save_total_limit=2,
    num_train_epochs=3,per_device_train_batch_size=64)
class ReadLineDS(object):
    def __init__(self,file,tokenizer):
        self.tokenizer=tokenizer
        with open(file,"r",encoding="utf-8") as r:
            self.lines=[s.strip() for s in r if s.strip()!=""]
    __len__=lambda self:len(self.lines)
    __getitem__=lambda self,i:self.tokenizer(self.lines[i],truncation=True,
        add_special_tokens=True,max_length=self.tokenizer.model_max_length-2)
trn=Trainer(args=arg,data_collator=DataCollatorForLanguageModeling(tkz),
    model=DebertaV2ForMaskedLM(cfg),train_dataset=ReadLineDS("train.txt",tkz))
trn.train()
trn.save_model("my-dir/deberta-uk")
tkz.save_pretrained("my-dir/deberta-uk")

```

図 91: ウクライナ語ミニ DeBERTa モデルの作成

```

s,d="my-dir/deberta-uk","my-dir/deberta-uk-upos"
!python -m esupar.train {s} {d} 32 /tmp train.upos
!python -m esupar.train {d} {d} 32 /// train.conllu dev.conllu test.conllu

```

図 92: ウクライナ語向け品詞付与・係り受け解析ミニモデルの製作

```

!pip install esupar
import esupar
nlp=esupar.load("my-dir/deberta-uk-upos")
doc=nlp("Не скупись на втіху їй і ласку любий брате")
import deplacy
deplacy.serve(doc,port=None)

```

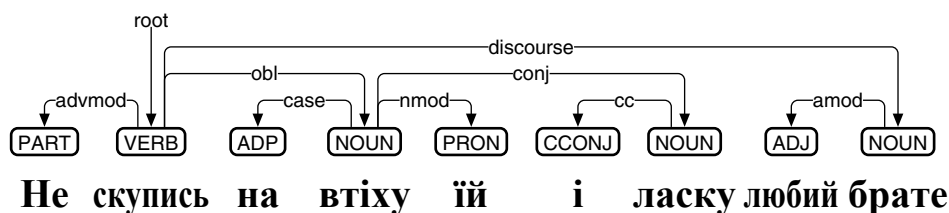


図 93: ウクライナ語向け品詞付与・係り受け解析ミニモデルのテスト

12 ベラルーシ語編

12.1 ベラルーシ語 UD における品詞付与と係り受け解析

deberta-base-belarusian-ud-goeswith^[102]は、隣接行列型係り受け解析アルゴリズムを、事前学習モデルで実装しています。Google Colaboratory 上の Transformers^[26]を用いて、deberta-base-belarusian-ud-goeswith でベラルーシ語の例文「Насамрэч інтэрнэт стварае міфы тых у чыіх руках інтэрнэт」を品詞付与・係り受け解析するプログラムと、その結果を図 94 に示します。解析結果を見ていくことにしましょう。

```
!pip install deplacy transformers
from transformers import pipeline
mdl="KoichiYasuoka/deberta-base-belarusian-ud-goeswith"
arg={"trust_remote_code":True,"aggregation_strategy":"simple"}
nlp=pipeline("universal-dependencies",mdl,**arg)
doc=nlp("Насамрэч інтэрнэт стварае міфы тых у чыіх руках інтэрнэт")
import deplacy
deplacy.serve(doc,port=None)
```

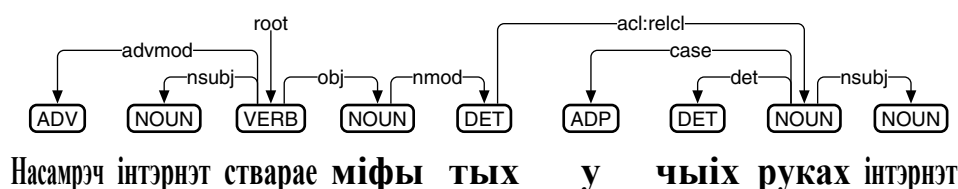


図 94: deberta-base-belarusian-ud-goeswith によるベラルーシ語の係り受け解析

2つの「інтэрнэт」は、前者が「стварае」の主語、後者が「у чыіх руках інтэрнэт」というコピュラ節の主語であり、いずれも *nsubj* で繋がられています。コピュラ節「у чыіх руках інтэрнэт」の内部では、前置詞「у」が *case* で、関係代名詞「чыіх」が *det* で繋がれており、コピュラ節全体が *acl:relcl* で「тых」を修飾しています。「тых」は「міфы」を修飾していて、*nmod* で繋がられています。「міфы」は「стварае」の目的語で、*obj* で繋がれています。「Насамрэч」は *advmod* で繋がれています。

なお、ベラルーシ語の係り受けタグには、表 4 に加え、受動態を表す *nsubj:pass* と *aux:pass*、受動態の動作主を表す *obl:agent*、関係代名詞による連体修飾節を表す *acl:relcl*、数詞と名詞の格変化の関係を表す *nummod:gov* と *nummod:entity*、コピュラ節に対する副詞的修飾語を表す *advmod:discourse*、氏名を表す *flat:name*、ベラルーシ語以外を表す *flat:foreign* も用いられます^[103]。

12.2 ベラルーシ語ミニ DeBERTa モデルの製作

Google Colaboratory 上の esupar^[47]で、ベラルーシ語ミニ DeBERTa モデルを製作してみましよう。手順を図 95～98 に示します。

^[102]<https://huggingface.co/KoichiYasuoka/deberta-base-belarusian-ud-goeswith>

^[103]<https://universaldependencies.org/be/#relation-subtypes>

最初に、訓練のための文章を準備しましょう (図 95)。train.txt の文章の材料として、ここでは UD_Belarusian-HSE^[104]を元にしてありますが、好みに応じ、ベラルーシ語の他の文章を増量するのもいいでしょう。なお、各 CoNLL-U ファイルは、品詞付与・係り受け解析の学習にも用います。

```
!pip install esupar accelerate
import os
url="https://github.com/UniversalDependencies/UD_Belarusian-HSE"
d=os.path.basename(url)
!test -d {d} || git clone --depth=1 {url}
!for F in train dev test ; do cp {d}/*-$$F.conllu $$F.conllu ; done
!sed -n "s/^# text = //p" *.conllu > train.txt
```

図 95: ベラルーシ語ミニモデル製作のための準備

続けて、ベラルーシ語向けトークナイザを、my-dir/tokenizer-be に作成しましょう。図 96 では、Transformers の BertTokenizerFast を使って、ベラルーシ語向けトークナイザを作成しています。語彙数(vocab_size)を 8000 に絞っていますが、もう少し大きくした方がいいかもしれません。

```
from transformers import BertTokenizerFast
from tokenizers import BertWordPieceTokenizer
s=["[CLS]","[PAD]","[SEP]","[UNK]","[MASK]"]
wpt=BertWordPieceTokenizer(strip_accents=False)
wpt.train(files=["train.txt"],vocab_size=8000,special_tokens=s)
wpt.save_model(".")
tkz=BertTokenizerFast(vocab_file="vocab.txt",strip_accents=False,never_split=s)
tkz.save_pretrained("my-dir/tokenizer-be")
```

図 96: ベラルーシ語向けトークナイザの作成

次に train.txt の文章から、DeBERTa モデルを my-dir/deberta-be に作成しましょう。図 97 では、入出力幅 128 トークン・入出力ベクトル 256 次元・深さ 12 層・アテンションヘッド 4 個・中間ベクトル 768 次元という、やや小さめのモデルにしています。

最後に、my-dir/deberta-be と UD_Belarusian-HSE から、品詞付与・係り受け解析モデル my-dir/deberta-be-upos を製作しましょう (図 98)。うまく my-dir/deberta-be-upos が製作できたら、esupar でテストしてみましょう (図 99)。

^[104]https://github.com/UniversalDependencies/UD_Belarusian-HSE

```

from transformers import (AutoTokenizer,DebertaV2Config,DebertaV2ForMaskedLM,
    DataCollatorForLanguageModeling,TrainingArguments,Trainer)
tkz=AutoTokenizer.from_pretrained("my-dir/tokenizer-be",model_max_length=128)
t=tkz.convert_tokens_to_ids(["[CLS]","[PAD]","[SEP]","[UNK]","[MASK]"])
cfg=DebertaV2Config(hidden_size=256,num_hidden_layers=12,num_attention_heads=4,
    intermediate_size=768,relative_attention=True,position_biased_input=False,
    pos_att_type=["p2c","c2p"],max_position_embeddings=tkz.model_max_length,
    vocab_size=len(tkz),tokenizer_class=type(tkz).__name__,
    bos_token_id=t[0],pad_token_id=t[1],eos_token_id=t[2])
arg=TrainingArguments(report_to="none",output_dir="/tmp",save_total_limit=2,
    num_train_epochs=3,per_device_train_batch_size=64)
class ReadLineDS(object):
    def __init__(self,file,tokenizer):
        self.tokenizer=tokenizer
        with open(file,"r",encoding="utf-8") as r:
            self.lines=[s.strip() for s in r if s.strip()!=""]
    __len__=lambda self:len(self.lines)
    __getitem__=lambda self,i:self.tokenizer(self.lines[i],truncation=True,
        add_special_tokens=True,max_length=self.tokenizer.model_max_length-2)
trn=Trainer(args=arg,data_collator=DataCollatorForLanguageModeling(tkz),
    model=DebertaV2ForMaskedLM(cfg),train_dataset=ReadLineDS("train.txt",tkz))
trn.train()
trn.save_model("my-dir/deberta-be")
tkz.save_pretrained("my-dir/deberta-be")

```

図 97: ベラルーシ語ミニ DeBERTa モデルの作成

```
!python -m esupar.train my-dir/deberta-be my-dir/deberta-be-upos .
```

図 98: ベラルーシ語向け品詞付与・係り受け解析ミニモデルの製作

```

!pip install esupar
import esupar
nlp=esupar.load("my-dir/deberta-be-upos")
doc=nlp("Насамрэч інтэрнэт стварае міфы тых у чыіх руках інтэрнэт")
import deplacy
deplacy.serve(doc,port=None)

```

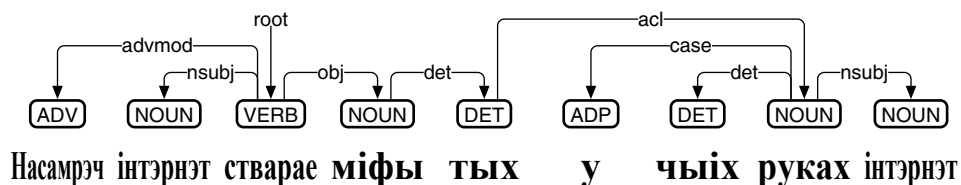


図 99: ベラルーシ語向け品詞付与・係り受け解析ミニモデルのテスト

13 セルビア語編

13.1 セルビア語 UD における品詞付与と係り受け解析

esupar^[47]は、隣接行列型係り受け解析アルゴリズムを、事前学習モデルで実装しています。Google Colaboratory 上の esupar を用いて、gpt2-small-serbian-upos^[105]でセルビア語の例文「Моја би мати знала гибати гибаницу да има сира и масла」を品詞付与・係り受け解析するプログラムと、その結果を図 100 に示します。解析結果を見ていくことにしましょう。

```
!pip install esupar
import esupar
nlp=esupar.load("KoichiYasuoka/gpt2-small-serbian-upos")
doc=nlp("Моја би мати знала гибати гибаницу да има сира и масла")
import deplacy
deplacy.serve(doc,port=None)
```

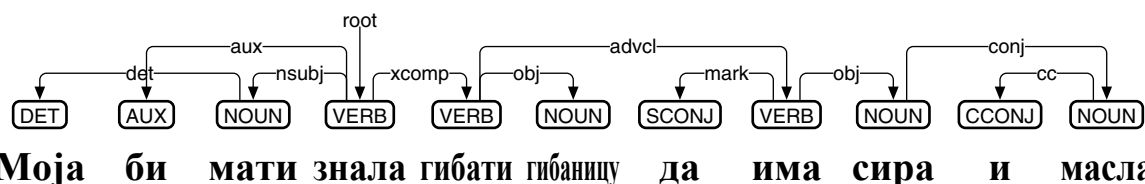


図 100: gpt2-small-serbian-upos によるセルビア語 (キリル文字) の係り受け解析

図 100 の例文は、主節「Моја би мати знала гибати гибаницу」と条件節「да има сира и масла」から構成されており、advcl で繋がられています。動詞「знала」の主語「мати」は、同時に「гибати」の主語でもあることから、「мати」が nsubj で、「гибати」が xcomp で繋がられています。助動詞「би」は aux で繋がられています。所有代名詞「Моја」は「мати」を修飾しており、det のリンクが交差しています。「гибаницу」は「гибати」の目的語であり、obj で繋がられています。条件節では、「да」が mark で繋がられています。「има」の目的語は「сира」と「масла」の 2 つですが、「сира」の方が obj で繋がれていて、その先に conj で「масла」が繋がれており、そのまた先に cc で「и」が繋がられています。

セルビア語は、公式にはキリル文字で書かれるのですが、ラテン文字による表記も多く用いられており、UD_Serbian-SET^[106]などのセルビア語コーパスも、多くはラテン文字で書かれています。gpt2-small-serbian-upos は、ラテン文字で書かれたセルビア語も、キリル文字と同様に解析可能です。Google Colaboratory 上の esupar を用いて、セルビア語の例文「Moja bi mati znala gibati gibanicu da ima sira i masla」を品詞付与・係り受け解析するプログラムと、その結果を図 101 に示します。

なお、セルビア語 UD の係り受けタグには、表 4 に加え、数詞と名詞 (および代名詞) の格変化の関係を表す nummod:gov と det:numgov も用いられます^[107]。

^[105]<https://huggingface.co/KoichiYasuoka/gpt2-small-serbian-upos>

^[106]Tanja Samardžić, Mirjana Starović, Željko Agić, Nikola Ljubešić: Universal Dependencies for Serbian in Comparison with Croatian and Other Slavic Languages, Proceedings of the 6th Workshop on Balto-Slavic Natural Language Processing (April 2017), pp.39-44.

^[107]https://universaldependencies.org/treebanks/sr_set/#relations-overview

```
!pip install esupar
import esupar
nlp=esupar.load("KoichiYasuoka/gpt2-small-serbian-upos")
doc=nlp("Moja bi mati znala gibati gibanicu da ima sira i masla")
import deplacy
deplacy.serve(doc,port=None)
```

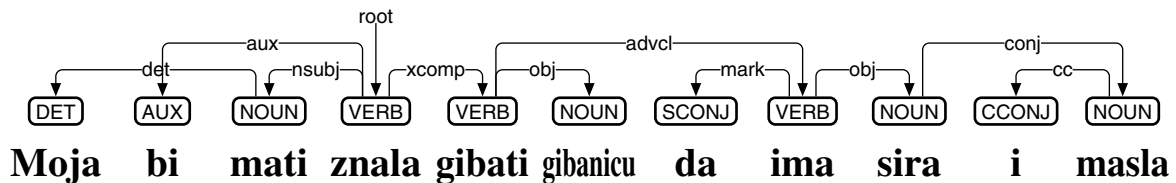


図 101: gpt2-small-serbian-upos によるセルビア語 (ラテン文字) の係り受け解析

13.2 セルビア語ミニ DeBERTa モデルの製作

Google Colaboratory 上の esupar で、セルビア語ミニ DeBERTa モデルを製作してみましょう。手順を図 102～106 に示します。srtools^[108]の助けを借りて、キリル文字とラテン文字の両方が使えるモデルを目指します。

最初に、訓練のための文章を準備しましょう (図 102)。train.txt の文章の材料として、ここでは srWaC^[109]の一部と、UD_Serbian-SET を元にしてはいますが、好みに応じ、セルビア語の他の文章を増量するのもいいでしょう。なお、UD_Serbian-SET の各 CoNLL-U ファイルは、品詞付与・係り受け解析の学習にも用います。

続けて、train.txt の各単語をキリル文字に変換しつつ、セルビア語向けトークナイザを、my-dir/tokenizer-sr に作成しましょう。図 103 では、Transformers^[26]の BertTokenizerFast を使って、セルビア語向けトークナイザを作成しています。語彙数(vocab_size)を 8000 に絞っていますが、もう少し大きくした方がいいかもしれません。

次に、train.txt の文章をキリル文字に変換しつつ、DeBERTa モデルを my-dir/deberta-sr に作成しましょう。図 104 では、入出力幅 128 トークン・入出力ベクトル 256 次元・深さ 12 層・アテンションヘッド 4 個・中間ベクトル 768 次元という、やや小さめのモデル

```
!pip install esupar srtools accelerate
import os
url="https://www.clarin.si/repository/xmlui/bitstream/handle/11356/1063"
!test -f srWaC1.1.06.xml.gz || curl -LO {url}/srWaC1.1.06.xml.gz
s='if($0~/^</){s="";if($1=="</s>")print""}else{printf("%s%s",s,$1);s=" "}}'
!zcat *.xml.gz | awk '{s}' | fgrep -v "&" > train.txt
url="https://github.com/UniversalDependencies/UD_Serbian-SET"
d=os.path.basename(url)
!test -d {d} || git clone --depth=1 {url}
!for F in train dev test ; do cp {d}/*-$$F.conllu $$F.conllu ; done
!sed -n "s/^# text = //p" *.conllu >> train.txt
```

図 102: セルビア語ミニモデル製作のための準備

^[108]<https://gitlab.com/andrejr/srtools>

^[109]<http://nlp.ffzg.hr/resources/corpora/srWaC>

```

from transformers import BertTokenizerFast
from tokenizers import BertWordPieceTokenizer
from srtools import latin_to_cyrillic
s=[" [CLS]", " [PAD]", " [SEP]", " [UNK]", " [MASK]"]
wpt=BertWordPieceTokenizer(strip_accents=False)
with open("train.txt","r",encoding="utf-8") as r:
    t=latin_to_cyrillic(r.read().strip()).split("\n")
wpt.train_from_iterator(iterator=t,vocab_size=8000,special_tokens=s)
wpt.save_model(".")
tkz=BertTokenizerFast(vocab_file="vocab.txt",strip_accents=False,never_split=s)
tkz.save_pretrained("my-dir/tokenizer-sr")

```

図 103: セルビア語向けトークナイザの作成

にしています。さらに、キリル文字のトークンに対応するラテン文字のトークンを追加し、新たなモデルを my-dir/deberta-sr-ext に保存しましょう (図 105)。

最後に、my-dir/deberta-sr-ext と UD_Serbian-SET から、品詞付与・係り受け解析モデル my-dir/deberta-sr-upos を製作しましょう (図 106)。うまく my-dir/deberta-sr-upos が製作できたら、esupar でテストしてみましょう。キリル文字とラテン文字の両方が使えるはずですが、図 107 の例では「гибати гибаницу」を固有名詞だと解釈してしまっており、図 100 とは異なる結果になっているようです。

```

from transformers import (AutoTokenizer,DebertaV2Config,DebertaV2ForMaskedLM,
    DataCollatorForLanguageModeling,TrainingArguments,Trainer)
tkz=AutoTokenizer.from_pretrained("my-dir/tokenizer-sr",model_max_length=128)
t=tkz.convert_tokens_to_ids([" [CLS]", " [PAD]", " [SEP]", " [UNK]", " [MASK]"])
cfg=DebertaV2Config(hidden_size=256,num_hidden_layers=12,num_attention_heads=4,
    intermediate_size=768,relative_attention=True,position_biased_input=False,
    pos_att_type=["p2c","c2p"],max_position_embeddings=tkz.model_max_length,
    vocab_size=len(tkz),tokenizer_class=type(tkz).__name__,
    bos_token_id=t[0],pad_token_id=t[1],eos_token_id=t[2])
arg=TrainingArguments(report_to="none",output_dir="/tmp",save_total_limit=2,
    num_train_epochs=3,per_device_train_batch_size=64)
class ReadLineCyr(object):
    def __init__(self,file,tokenizer):
        from srtools import latin_to_cyrillic
        self.tokenizer=tokenizer
        with open(file,"r",encoding="utf-8") as r:
            self.lines=[latin_to_cyrillic(s.strip()) for s in r if s.strip()!=""]
        __len__=lambda self:len(self.lines)
        __getitem__=lambda self,i:self.tokenizer(self.lines[i],truncation=True,
            add_special_tokens=True,max_length=self.tokenizer.model_max_length-2)
trn=Trainer(args=arg,data_collator=DataCollatorForLanguageModeling(tkz),
    model=DebertaV2ForMaskedLM(cfg),train_dataset=ReadLineCyr("train.txt",tkz))
trn.train()
trn.save_model("my-dir/deberta-sr")
tkz.save_pretrained("my-dir/deberta-sr")

```

図 104: セルビア語ミニ DeBERTa モデルの作成

```

import torch,srtools
from transformers import BertTokenizerFast,AutoModelForMaskedLM
tkz=BertTokenizerFast.from_pretrained("my-dir/deberta-sr")
mdl=AutoModelForMaskedLM.from_pretrained("my-dir/deberta-sr")
n,s,m=len(tkz),tkz.all_special_tokens,tkz.model_max_length
c=[t if t in s else srtools.cyrillic_to_latin(t)
   for t in tkz.convert_ids_to_tokens(range(n))]
t=[False]*n+[True]*tkz.add_tokens(c)
e=mdl.resize_token_embeddings(len(tkz))
with torch.no_grad():
    for i,j in enumerate(tkz.convert_tokens_to_ids(c)):
        if t[j]:
            e.weight[j,:]=e.weight[i,:]
            t[j]=False
mdl.set_input_embeddings(e)
mdl.save_pretrained("my-dir/deberta-sr-ext")
with open("vocab.txt","w",encoding="utf-8") as w:
    print("\n".join(tkz.convert_ids_to_tokens(range(len(tkz))))),file=w)
tkz=BertTokenizerFast(vocab_file="vocab.txt",strip_accents=False,never_split=s,
    model_max_length=m)
tkz.save_pretrained("my-dir/deberta-sr-ext")

```

図 105: セルビア語ミニ DeBERTa モデルの拡張 (ラテン文字追加)

```
!python -m esupar.train my-dir/deberta-sr-ext my-dir/deberta-sr-upos .
```

図 106: セルビア語向け品詞付与・係り受け解析ミニモデルの製作

```

!pip install esupar
import esupar
nlp=esupar.load("my-dir/deberta-sr-upos")
doc=nlp("Moja bi mati znala gibati gibаницу да има сира и масла")
import deplacy
deplacy.serve(doc,port=None)

```

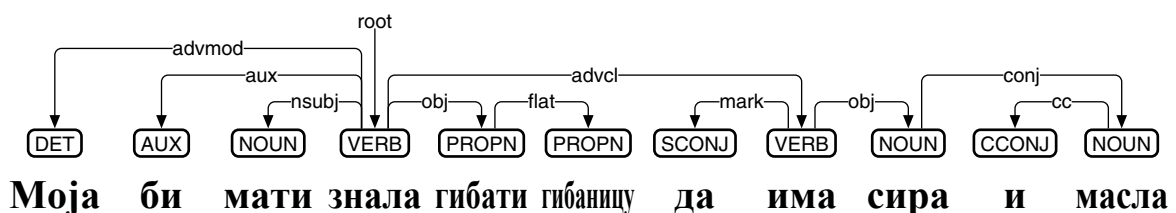


図 107: セルビア語向け品詞付与・係り受け解析ミニモデルのテスト

14 クロアチア語編

14.1 クロアチア語 UD における品詞付与と係り受け解析

Trankit^[64]は、XLM-RoBERTa をベースに、隣接行列型係り受け解析アルゴリズムを、様々な言語向けに実装しています。Google Colaboratory 上の Trankit を用いて、クロアチア語の例文「Moja bi mati znala gibati gibanicu da ima sira i masla」を品詞付与・係り受け解析するプログラムと、その結果を図 108 に示します。解析結果を見ていくことにしましょう。

```
!pip install git+https://github.com/KoichiYasuoka/trankit deplacy
import trankit
nlp=trankit.Pipeline("croatian")
doc=nlp("Moja bi mati znala gibati gibanicu da ima sira i masla")
import deplacy
deplacy.serve(doc,port=None)
```

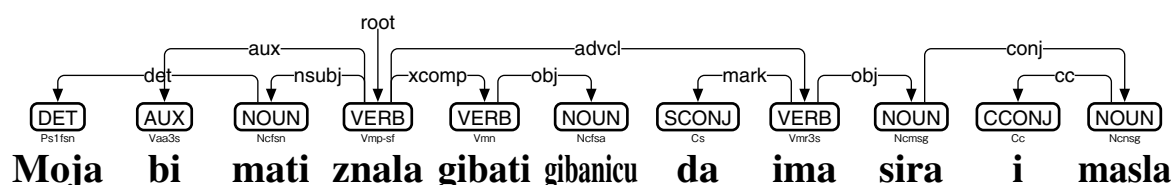


図 108: Trankit によるクロアチア語の係り受け解析

図 108 の例文は、主節「Moja bi mati znala gibati gibanicu」と条件節「da ima sira i masla」から構成されており、advcl で繋がられています。動詞「znala」の主語「mati」は、同時に「gibati」の主語でもあることから、「mati」が nsubj で、「gibati」が xcomp で繋がられています。助動詞「bi」は aux で繋がられています。所有代名詞「Moja」は「mati」を修飾しており、det のリンクが交差しています。「gibanicu」は「gibati」の目的語であり、obj で繋がられています。条件節では、「da」が mark で繋がられています。「ima」の目的語は「sira」と「masla」の 2 つですが、「sira」の方が obj で繋がれていて、その先に conj で「masla」が繋がれており、そのまた先に cc で「i」が繋がられています。

なお、クロアチア語 UD の係り受けタグには、表 4 に加え、強調の連用修飾語を表す advmod:emph、数詞と代名詞の格変化の関係を表す det:numgov、クロアチア語以外を表す flat:foreign も用いられます^[110]。

14.2 クロアチア語ミニ DeBERTa モデルの製作

Google Colaboratory 上の esupar^[47]で、クロアチア語ミニ DeBERTa モデルを製作してみましょう。手順を図 109～112 に示します。

最初に、訓練のための文章を準備しましょう (図 109)。train.txt の文章の材料として、ここでは UD_Croatian-SET^[111]と hr500k^[112]を元にしていますが、好みに応じ、クロアチア語

^[110]https://universaldependencies.org/treebanks/hr_set/#relations-overview

^[111]https://github.com/UniversalDependencies/UD_Croatian-SET

^[112]<https://huggingface.co/datasets/classla/hr500k>

```

!pip install esupar accelerate
import os
url="https://github.com/UniversalDependencies/UD_Croatian-SET"
d=os.path.basename(url)
!test -d {d} || git clone --depth=1 {url}
!for F in train dev test ; do cp {d}/*-$$F.conllu $$F.conllu ; done
url="https://huggingface.co/datasets/classla/hr500k"
d=os.path.basename(url)
!test -d {d} || ( git clone --depth=1 {url} ; cd {d} ; unzip data_ner.zip )
!cat {d}/data_ner/*.conllu *.conllu > train.upos
!sed -n "s/^# text = //p" train.upos > train.txt

```

図 109: クロアチア語ミニモデル製作のための準備

```

from transformers import BertTokenizerFast
from tokenizers import BertWordPieceTokenizer
s=["[CLS]","[PAD]","[SEP]","[UNK]","[MASK]"]
wpt=BertWordPieceTokenizer(lowercase=False,strip_accents=False)
wpt.train(files=["train.txt"],vocab_size=8000,special_tokens=s)
wpt.save_model(".")
tkz=BertTokenizerFast(vocab_file="vocab.txt",do_lowercase=False,
    strip_accents=False,never_split=s)
tkz.save_pretrained("my-dir/tokenizer-hr")

```

図 110: クロアチア語向けトークナイザの作成

の他の文章を増量するのもいいでしょう。なお、UD_Croatian-SET・hr500k の各 CoNLL-U ファイルは、train.upos にまとめ上げて、品詞付与の学習に用います。また、UD_Croatian-SET の各 CoNLL-U ファイルは、係り受け解析の学習にも用います。

続けて、クロアチア語向けトークナイザを、my-dir/tokenizer-hr に作成しましょう。図 110 では、Transformers^[26] の BertTokenizerFast を使って、クロアチア語向けトークナイザを作成しています。語彙数(vocab_size)を 8000 に絞っていますが、もう少し大きくした方がいいかもしれません。

次に train.txt の文章から、DeBERTa モデルを my-dir/deberta-hr に作成しましょう。図 111 では、入出力幅 128 トークン・入出力ベクトル 256 次元・深さ 12 層・アテンションヘッド 4 個・中間ベクトル 768 次元という、やや小さめのモデルにしています。

最後に、my-dir/deberta-hr と UD_Croatian-SET・hr500k から、品詞付与・係り受け解析モデル my-dir/deberta-hr-upos を製作しましょう(図 112)。うまく my-dir/deberta-hr-upos が製作できたら、esupar でテストしてみましょう。ただ、図 113 の例では「gibati gibanicu」を固有名詞だと解釈してしまっており、図 108 とは異なる結果になっているようです。

```

from transformers import (AutoTokenizer,DebertaV2Config,DebertaV2ForMaskedLM,
    DataCollatorForLanguageModeling,TrainingArguments,Trainer)
tkz=AutoTokenizer.from_pretrained("my-dir/tokenizer-hr",model_max_length=128)
t=tkz.convert_tokens_to_ids(["[CLS]","[PAD]","[SEP]","[UNK]","[MASK]"])
cfg=DebertaV2Config(hidden_size=256,num_hidden_layers=12,num_attention_heads=4,
    intermediate_size=768,relative_attention=True,position_biased_input=False,
    pos_att_type=["p2c","c2p"],max_position_embeddings=tkz.model_max_length,
    vocab_size=len(tkz),tokenizer_class=type(tkz).__name__,
    bos_token_id=t[0],pad_token_id=t[1],eos_token_id=t[2])
arg=TrainingArguments(report_to="none",output_dir="/tmp",save_total_limit=2,
    num_train_epochs=3,per_device_train_batch_size=64)
class ReadLineDS(object):
    def __init__(self,file,tokenizer):
        self.tokenizer=tokenizer
        with open(file,"r",encoding="utf-8") as r:
            self.lines=[s.strip() for s in r if s.strip()!=""]
    __len__=lambda self:len(self.lines)
    __getitem__=lambda self,i:self.tokenizer(self.lines[i],truncation=True,
        add_special_tokens=True,max_length=self.tokenizer.model_max_length-2)
trn=Trainer(args=arg,data_collator=DataCollatorForLanguageModeling(tkz),
    model=DebertaV2ForMaskedLM(cfg),train_dataset=ReadLineDS("train.txt",tkz))
trn.train()
trn.save_model("my-dir/deberta-hr")
tkz.save_pretrained("my-dir/deberta-hr")

```

図 111: クロアチア語ミニ DeBERTa モデルの作成

```

s,d="my-dir/deberta-hr","my-dir/deberta-hr-upos"
!python -m esupar.train {s} {d} 32 /tmp train.upos
!python -m esupar.train {d} {d} 32 /// train.conllu dev.conllu test.conllu

```

図 112: クロアチア語向け品詞付与・係り受け解析ミニモデルの製作

```

!pip install esupar
import esupar
nlp=esupar.load("my-dir/deberta-hr-upos")
doc=nlp("Moja bi mati znala gibati gibanicu da ima sira i masla")
import deplacy
deplacy.serve(doc,port=None)

```

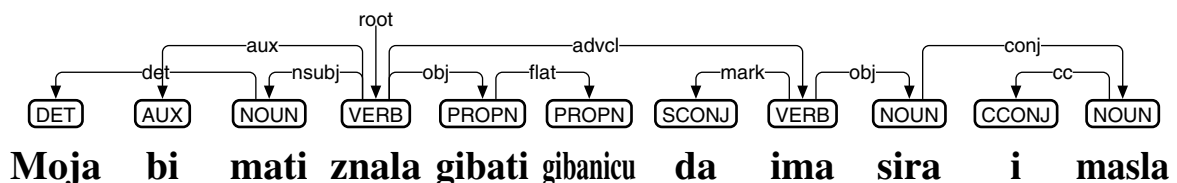


図 113: クロアチア語向け品詞付与・係り受け解析ミニモデルのテスト

15 コプト語編

15.1 コプト語UDにおける品詞付与

書写言語としてのコプト語は、単語と単語の間に区切りがありません。空白や改行が入っている場合が多いのですが、単語の区切りとは限らないのです。事前学習モデルでコプト語を扱う際には、各トークンを単語と合致させるのが理想です。しかし、それは現実には困難なため、各トークンは単語より短めにトークナイズした上で、系列ラベリング(UPOS 付与)によって、単語をまとめ上げる手法が考えられます。すなわち、トークンが単語と合致した場合には UPOS をラベルとして付与し、そうでない場合には、最初のトークンに「B-」を付けた UPOS を、それ以降のトークンに「I-」を付けた UPOS を、それぞれ付与するのです。

```
!pip install deplacy transformers
class SeqL(object):
    def __init__(self,bert):
        from transformers import pipeline
        self.tagger=pipeline(task="ner",model=bert)
    def __call__(self,text):
        w=[(t["start"],t["end"],t["entity"],t["word"]) for t in self.tagger(text)]
        u,z="# text = "+text.replace("\n"," ")+"\n","_"
        for i,(s,e,p,t) in enumerate(w,1):
            m,q=z if i<len(w) and e<w[i][0] else "SpaceAfter=No",p.split("|")
            f=z if p.find("=")<0 else "|".join(t for t in q if t.find("=")>0)
            u+="\t".join([str(i),t.strip(),z,q[0].replace(".", ""),z,f,z,z,z,m])+ "\n"
        return u+"\n"
nlp=SeqL("KoichiYasuoka/roberta-base-coptic-upos")
doc=nlp("μοοϣεζωσϣηρεμπογοειν")
import deplacy
deplacy.serve(doc,port=None)
```

(B-VERB)	(I-VERB)	(I-VERB)	(SCONJ)	(NOUN)	(ADP)	(DET)	(NOUN)
μ	οοϣ	ε	ζωσ	ϣηρε	μ	π	ογοειν

図 114: roberta-base-coptic-upos による UPOS 付与

Google Colaboratory 上の Transformers^[26]を用いて、roberta-base-coptic-upos^[113]でコプト語の例文「μοοϣεζωσϣηρεμπογοειν」に UPOS を付与するプログラムと、その結果を図 114 に示します。「ζωσ」「ϣηρε」「μ」「π」「ογοειν」の各単語では、各トークンがそのまま単語と合致しており、それぞれ UPOS 品詞 SCONJ・NOUN・ADP・DET・NOUN が付与されています。「μοοϣε」は、3つのトークンに泣き別れてしまっており、それぞれに B-VERB・I-VERB・I-VERB が付与されています。

^[113]<https://huggingface.co/KoichiYasuoka/roberta-base-coptic-upos>

15.2 コプト語 UD における係り受け解析

esupar^[47]は、隣接行列型係り受け解析アルゴリズムを、事前学習モデルで実装しています。Google Colaboratory 上の esupar を用いて、roberta-base-coptic-upos でコプト語の例文「**ἡ ἀρετὴ ἐκ τῆς ἀρετῆς ἐκτελεῖται**」を係り受け解析するプログラムと、その結果を図 115 に示します。解析結果を見ていくことにしましょう。

```
!pip install esupar
import esupar
nlp=esupar.load("KoichiYasuoka/roberta-base-coptic-upos")
doc=nlp("ἡ ἀρετὴ ἐκ τῆς ἀρετῆς ἐκτελεῖται")
import deplacy
deplacy.serve(doc,port=None)
```

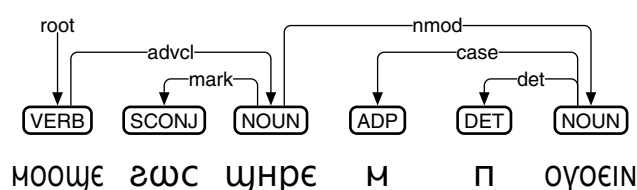


図 115: roberta-base-coptic-upos によるコプト語の係り受け解析

図 115 の例文は、主節「**ἡ ἀρετὴ ἐκτελεῖται**」と従属節「**ἐκ τῆς ἀρετῆς ἐκτελεῖται**」から構成されており、advcl で繋がられています。主節は動詞「**ἐκτελεῖται**」のみです。従属節はコピュラ文(ただし主語も繫辞もない)で、「**ἐκ**」が mark で繋がられています。「**ἀρετῆς**」は「**ἀρετὴ**」を修飾していて、nmod で繋がられています。「**ἀρετῆς**」は、前置詞「**ἐκ**」と冠詞「**ἡ**」と名詞「**ἀρετὴ**」から成っていて、それぞれ case と det で繋がられています。

なお、コプト語 UD の係り受けタグには、表 4 に加え、斜格補語による連用修飾を表す obl:nmod、関係転換詞による連体修飾節を表す acl:relcl も用いられます^[114]。

15.3 コプト語ミニ DeBERTa モデルの製作

Google Colaboratory 上の esupar で、コプト語ミニ DeBERTa モデルを製作してみましょう。手順を図 116～119 に示します。

最初に、訓練のための文章を準備しましょう(図 116)。train.txt の文章の材料として、ここでは Coptic Scriptorium Corpora^[115]と UD_Coptic-Scriptorium^[116]を元にしていますが、好みに応じ、コプト語の他の文章を増量するのもいいでしょう。train.txt と同時に、各単語を空白で区切った token.txt も準備しており、トークナイザの作成に用います。Coptic Scriptorium Corpora の各 CoNLL-U ファイルは、train.upos にまとめ上げて、品詞付与の学習に用います。また、UD_Coptic-Scriptorium の各 CoNLL-U ファイルは、係り受け解析の学習に用います。

続けて、token.txt の各単語から、コプト語向けトークナイザを、my-dir/tokenizer-cop に作成しましょう。図 117 では、Unigram トークナイザ (SentencePiece^[88]の改造版) を

^[114]<https://universaldependencies.org/cop/dep>

^[115]<http://data.copticscriptorium.org/index/corpus>

^[116]https://github.com/UniversalDependencies/UD_Coptic-Scriptorium

```

!pip install esupar spacy-alignments accelerate
import os
url="https://github.com/UniversalDependencies/UD_Coptic-Scriptorium"
d=os.path.basename(url)
!test -d {d} || git clone --depth=1 {url}
!for F in train dev test ; do cp {d}/*-$$F.conllu $$F.conllu ; done
url="https://github.com/CopticScriptorium/corpora"
d=os.path.basename(url)
!test -d {d} || git clone --depth=1 {url}
s='''{if(NF==10&&$1~/^[1-9][0-9]*$/)}printf($1>1?" %s":"%s",$2);if(NF==0){print;
  if(u!~/\n[^\n1-9]/)}print u>"train.upos";u=""}else u=u$0"\n"}'''
!nawk -F\t' '{s}' {d}/**/*.conllu | tee token.txt > train.txt
!sed -n "s/^# text = //p" train.upos *.conllu >> train.txt

```

図 116: コプト語ミニモデル製作のための準備

```

from transformers import DebertaV2TokenizerFast
from tokenizers import (Tokenizer,models,pre_tokenizers,normalizers,processors,
  decoders,trainers)
s=[" [CLS] ", " [PAD] ", " [SEP] ", " [UNK] ", " [MASK] "]
spt=Tokenizer(models.Unigram())
spt.pre_tokenizer=pre_tokenizers.Sequence([pre_tokenizers.Whitespace(),
  pre_tokenizers.Punctuation()])
spt.normalizer=normalizers.Sequence([normalizers.Nmt(),normalizers.NFD(),
  normalizers.StripAccents()])
spt.post_processor=processors.TemplateProcessing(single=" [CLS] $A [SEP] ",
  pair=" [CLS] $A [SEP] $B:1 [SEP]:1",special_tokens=[(" [CLS] ",0),(" [SEP] ",2)])
spt.decoder=decoders.WordPiece(prefix="",cleanup=True)
spt.train(trainer=trainers.UnigramTrainer(vocab_size=3000,max_piece_length=8,
  special_tokens=s,unk_token=" [UNK] ",n_sub_iterations=2),files=["token.txt"])
spt.save("tokenizer.json")
tkz=DebertaV2TokenizerFast(tokenizer_file="tokenizer.json",split_by_punct=True,
  do_lower_case=False,keep_accents=False,vocab_file="/dev/null")
tkz.save_pretrained("my-dir/tokenizer-cop")

```

図 117: コプト語向けトークナイザの作成

token.txt に適用し、それを DebertaV2TokenizerFast という Transformers のトークナイザに組み上げて、コプト語向けトークナイザを作成しています。語彙数を 3000 に絞った上で、各トークンの文字数を最大 8 文字 (アクセント記号は除去) としています。

次に train.txt の文章から、DeBERTa モデルを my-dir/deberta-cop に作成しましょう。図 118 では、入出力幅 128 トークン・入出力ベクトル 256 次元・深さ 12 層・アテンションヘッド 4 個・中間ベクトル 768 次元という、やや小さめのモデルにしていますが、GPU でも 15 分程度かかってしまいます。

最後に、my-dir/deberta-cop と各 ConLL-U ファイルから、品詞付与・係り受け解析モデル my-dir/deberta-cop-upos を製作しましょう (図 119)。うまく my-dir/deberta-cop-upos ができたら、esupar でテストしてみましょう。ただ、図 120 の例では、「ϥⲙⲏⲣⲉ」を「ⲙⲟⲟⲩⲥⲉ」の斜格補語 (obl) だと解釈してしまっており、図 115 とは異なる結果になっているようです。

```

from transformers import (AutoTokenizer,DebertaV2Config,DebertaV2ForMaskedLM,
    DataCollatorForLanguageModeling,TrainingArguments,Trainer)
tkz=AutoTokenizer.from_pretrained("my-dir/tokenizer-cop",model_max_length=128)
t=tkz.convert_tokens_to_ids(["[CLS]","[PAD]","[SEP]","[UNK]","[MASK]"])
cfg=DebertaV2Config(hidden_size=256,num_hidden_layers=12,num_attention_heads=4,
    intermediate_size=768,relative_attention=True,position_biased_input=False,
    pos_att_type=["p2c","c2p"],max_position_embeddings=tkz.model_max_length,
    vocab_size=len(tkz),tokenizer_class=type(tkz).__name__,
    bos_token_id=t[0],pad_token_id=t[1],eos_token_id=t[2])
arg=TrainingArguments(report_to="none",output_dir="/tmp",save_total_limit=2,
    num_train_epochs=3,per_device_train_batch_size=64)
class ReadLineDS(object):
    def __init__(self,file,tokenizer):
        self.tokenizer=tokenizer
        with open(file,"r",encoding="utf-8") as r:
            self.lines=[s.strip() for s in r if s.strip()!=""]
    __len__=lambda self:len(self.lines)
    __getitem__=lambda self,i:self.tokenizer(self.lines[i],truncation=True,
        add_special_tokens=True,max_length=self.tokenizer.model_max_length-2)
trn=Trainer(args=arg,data_collator=DataCollatorForLanguageModeling(tkz),
    model=DebertaV2ForMaskedLM(cfg),train_dataset=ReadLineDS("train.txt",tkz))
trn.train()
trn.save_model("my-dir/deberta-cop")
tkz.save_pretrained("my-dir/deberta-cop")

```

図 118: コプト語ミニ DeBERTa モデルの作成

```

s,d="my-dir/deberta-cop","my-dir/deberta-cop-upos"
!python -m esupar.train {s} {d} 32 /tmp train.upos
!python -m esupar.train {d} {d} 32 /// train.conllu dev.conllu test.conllu

```

図 119: コプト語向け品詞付与・係り受け解析ミニモデルの製作

```

!pip install esupar
import esupar
nlp=esupar.load("my-dir/deberta-cop-upos")
doc=nlp("ⲙⲟⲟⲩⲉⲗⲱⲥⲩⲏⲣⲉⲙⲡⲟⲩⲟⲩⲓⲛ")
import deplacy
deplacy.serve(doc,port=None)

```

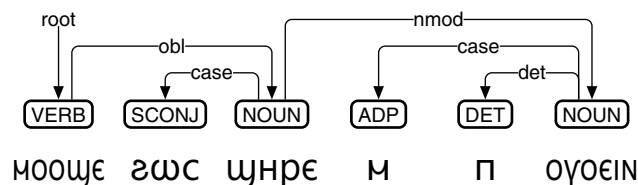


図 120: コプト語向け品詞付与・係り受け解析ミニモデルのテスト

16 スペイン語編

16.1 スペイン語 UD における品詞付与と係り受け解析

spacy-transformers^[73]は、BERT など事前学習モデルによる状態遷移型係り受け解析アルゴリズムを実装しており、スペイン語向けには、bert-base-spanish-wwm-cased^[117]を元にした es_dep_news_trf を公開^[118]しています。Google Colaboratory 上の spacy-transformers を用いて、スペイン語の例文「Deberías aclararlo al castellano y al catalán」を品詞付与・係り受け解析するプログラムと、その結果を図 121 に示します。解析結果を見ていくことにしましょう。

```
!pip install spacy-transformers deplacy
!python -m spacy download es_dep_news_trf
import spacy
nlp=spacy.load("es_dep_news_trf")
doc=nlp("Deberías aclararlo al castellano y al catalán")
import deplacy
deplacy.serve(doc,port=None)
```

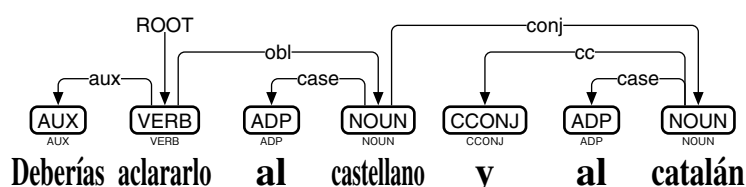


図 121: spacy-transformers によるスペイン語の係り受け解析

動詞「aclarar」は目的語「lo」が付いた形となっており、助動詞「Deberías」と aux で繋がっています。斜格補語は「castellano」と「catalán」の2つですが、「castellano」の方が obl で繋がれていて、その先に conj で「catalán」が繋がれており、そのまた先に cc で「y」が繋がれていると同時に、「castellano」と「catalán」にそれぞれ「al」が case で繋がれています。なお、spacy-transformers は「al」を1語だとみなしていますが、スペイン語 UD では「a」「el」の2語に分離した上で扱うことになっており、縮合語の解析が不十分です。

UDPipe 2 WebAPI は、UDPipe 2 の係り受け解析モデル^[67]を、WebAPI の形で公開しています。内部的には、隣接行列型係り受け解析アルゴリズムを使っています。Google Colaboratory から UDPipe 2 WebAPI の spanish-ancora にアクセスするプログラムと、スペイン語の例文「Deberías aclararlo al castellano y al catalán」を品詞付与・係り受け解析した結果を、図 122 に示します。「al」が内部的には2語として扱われており、「a」が前置詞、「el」が限定詞(冠詞)として、それぞれ case・det で繋がれています。また、「aclararlo」が2語として扱われており、「aclarar」に「lo」が obj で繋がれています。「aclarar」と「castellano」は obl:arg で繋がれていて、斜格補語でありながら項(argumento)だとみなされています。

^[117]José Cañete, Gabriel Chaperon, Rodrigo Fuentes, Jorge Pérez: Spanish pre-trained BERT model and evaluation data, ICLR 2020 (April 26, 2020).

^[118]https://huggingface.co/spacy/es_dep_news_trf


```

!pip install esupar spacy-alignments accelerate
import os
url="https://github.com/UniversalDependencies/UD_Spanish-AnCora"
!test -d {os.path.basename(url)} || git clone --depth=1 {url}
url="https://github.com/UniversalDependencies/UD_Spanish-GSD"
!test -d {os.path.basename(url)} || git clone --depth=1 {url}
!for F in train dev test ; do cat UD_Spanish-*/*-$$F.conllu > $$F.conllu ; done
!sed -n "s/^# text = //p" *.conllu > train.txt

```

図 123: スペイン語ミニモデル製作のための準備

```

from transformers import BertTokenizerFast
from tokenizers import BertWordPieceTokenizer
s=["[CLS]","[PAD]","[SEP]","[UNK]","[MASK]"]
wpt=BertWordPieceTokenizer(strip_accents=False)
wpt.train(files=["train.txt"],vocab_size=30000,special_tokens=s)
wpt.save_model(".")
tkz=BertTokenizerFast(vocab_file="vocab.txt",strip_accents=False,never_split=s)
tkz.save_pretrained("my-dir/tokenizer-es")

```

図 124: スペイン語向けトークナイザの作成

```

from transformers import (AutoTokenizer,DebertaV2Config,DebertaV2ForMaskedLM,
    DataCollatorForLanguageModeling,TrainingArguments,Trainer)
tkz=AutoTokenizer.from_pretrained("my-dir/tokenizer-es",model_max_length=128)
t=tkz.convert_tokens_to_ids(["[CLS]","[PAD]","[SEP]","[UNK]","[MASK]"])
cfg=DebertaV2Config(hidden_size=256,num_hidden_layers=12,num_attention_heads=4,
    intermediate_size=768,relative_attention=True,position_biased_input=False,
    pos_att_type=["p2c","c2p"],max_position_embeddings=tkz.model_max_length,
    vocab_size=len(tkz),tokenizer_class=type(tkz).__name__,
    bos_token_id=t[0],pad_token_id=t[1],eos_token_id=t[2])
arg=TrainingArguments(report_to="none",output_dir="/tmp",save_total_limit=2,
    num_train_epochs=3,per_device_train_batch_size=64)
class ReadLineDS(object):
    def __init__(self,file,tokenizer):
        self.tokenizer=tokenizer
        with open(file,"r",encoding="utf-8") as r:
            self.lines=[s.strip() for s in r if s.strip()!=""]
    __len__=lambda self:len(self.lines)
    __getitem__=lambda self,i:self.tokenizer(self.lines[i],truncation=True,
        add_special_tokens=True,max_length=self.tokenizer.model_max_length-2)
trn=Trainer(args=arg,data_collator=DataCollatorForLanguageModeling(tkz),
    model=DebertaV2ForMaskedLM(cfg),train_dataset=ReadLineDS("train.txt",tkz))
trn.train()
trn.save_model("my-dir/deberta-es")
tkz.save_pretrained("my-dir/deberta-es")

```

図 125: スペイン語ミニ DeBERTa モデルの作成

最後に、my-dir/deberta-es と UD_Spanish-Ancora・UD_Spanish-GSD から、品詞付与・係り受け解析モデル my-dir/deberta-es-upos を製作しましょう (図 126)。うまく my-dir/deberta-es-upos が製作できたら、esupar でテストしてみましょう。図 127 の例では「aclararlo」が 1 語とみなされており、図 122 とは異なる結果になっているようです。

```
!python -m esupar.train my-dir/deberta-es my-dir/deberta-es-upos .
```

図 126: スペイン語向け品詞付与・係り受け解析ミニモデルの製作

```
!pip install esupar
import esupar
nlp=esupar.load("my-dir/deberta-es-upos")
doc=nlp("Deberías aclararlo al castellano y al catalán")
import deplacy
deplacy.serve(doc,port=None)
```

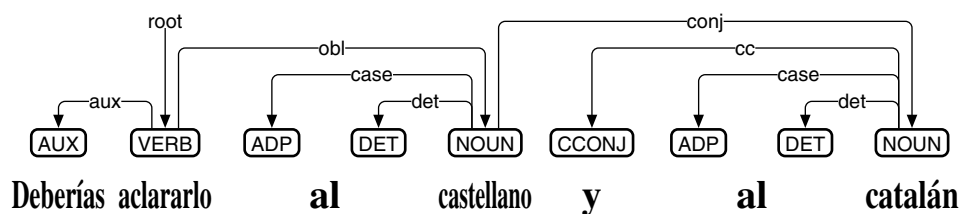


図 127: スペイン語向け品詞付与・係り受け解析ミニモデルのテスト

17 カタルーニャ語編

17.1 カタルーニャ語 UD における品詞付与と係り受け解析

Trankit^[64]は、XLM-RoBERTa をベースに、隣接行列型係り受け解析アルゴリズムを、様々な言語向けに実装しています。Google Colaboratory 上の Trankit を用いて、カタルーニャ語の例文「Ho hauries d'aclarir al català i al castellà」を品詞付与・係り受け解析するプログラムと、その結果を図 128 に示します。解析結果を見ていくことにしましょう。

```
!pip install git+https://github.com/KoichiYasuoka/trankit deplacy
import trankit
nlp=trankit.Pipeline("catalan")
doc=nlp("Ho hauries d'aclarir al català i al castellà")
import deplacy
deplacy.serve(doc,port=None)
```

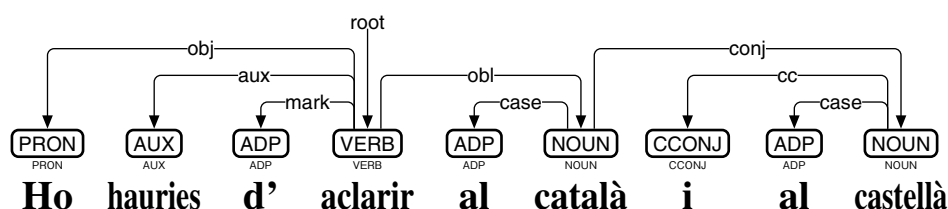


図 128: Trankit によるカタルーニャ語の係り受け解析

動詞「aclarir」には、目的語「Ho」が **obj** で繋がれており、さらに助動詞「hauries」と「d\'」が **aux** と **mark** で繋がれています。斜格補語は「català」と「castellà」の2つですが、「català」の方が **obl** で繋がれていて、その先に **conj** で「castellà」が繋がれており、そのまた先に **cc** で「i」が繋がれていると同時に、「català」と「castellà」にそれぞれ「al」が **case** で繋がれています。なお、Trankit は「al」を1語だとみなしていますが、カタルーニャ語 UD では「a」「el」の2語に分離した上で扱うことになっており、縮合語の解析が不十分です。

UDPipe 2 WebAPI は、UDPipe 2 の係り受け解析モデル^[67]を、WebAPI の形で公開しています。内部的には、隣接行列型係り受け解析アルゴリズムを使っています。Google Colaboratory から UDPipe 2 WebAPI の catalan-ancora にアクセスするプログラムと、カタルーニャ語の例文「Ho hauries d'aclarir al català i al castellà」を品詞付与・係り受け解析した結果を、図 129 に示します。「al」が内部的には2語として扱われており、「a」が前置詞、「el」が限定詞(冠詞)として、それぞれ **case**・**det** で繋がれています。

なお、カタルーニャ語 UD の係り受けタグには、表 4 に加え、受動態を表す **nsubj:pass**・**csubj:pass**・**expl:pass**・**aux:pass**、関係代名詞による連体修飾節を表す **acl:relcl**、再帰動詞における再帰代名詞を表す **expl:pv** も用いられます^[122]。

^[122]<https://universaldependencies.org/ca/#syntax>

```

!pip install deplacy
class UDPipe2WebAPI(object):
    def __init__(self, lang):
        self.url="https://lindat.mff.cuni.cz/services/udpipe/api/process"
        self.params={"model":lang,"tokenizer":"","tagger":"","parser":""}
    def __call__(self, text):
        import urllib.parse, urllib.request, json
        self.params["data"]=urllib.parse.quote(text)
        u=self.url+"?"+"&".join(k+"="+v for k,v in self.params.items())
        with urllib.request.urlopen(u) as r:
            return json.loads(r.read())["result"]
nlp=UDPipe2WebAPI("catalan-ancora")
doc=nlp("Ho hauries d'aclarir al català i al castellà")
import deplacy
deplacy.serve(doc, port=None)

```

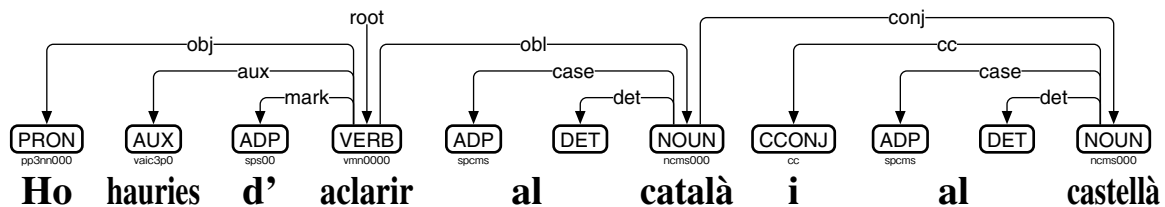


図 129: UDPipe 2 WebAPI によるカタルーニャ語の係り受け解析

17.2 カタルーニャ語ミニ DeBERTa モデルの製作

Google Colaboratory 上の esupar^[47]で、カタルーニャ語ミニ DeBERTa モデルを製作してみましょう。手順を図 130～133 に示します。

最初に、訓練のための文章を準備しましょう (図 130)。train.txt の文章の材料として、ここでは UD_Catalan-Ancora^[123]と Softcatalà コーパス^[124]を元にしてはいますが、好みに応じ、カタルーニャ語の他の文章を増量するのもいいでしょう。なお、UD_Catalan-Ancora の各 CoNLL-U ファイルは、品詞付与・係り受け解析の学習にも用います。

続けて、カタルーニャ語向けトークナイザを、my-dir/tokenizer-ca に作成しましょう。図 131 では、Transformers^[26]の BertTokenizerFast を使って、カタルーニャ語向けトークナイザを作成しています。

次に train.txt の文章から、DeBERTa モデルを my-dir/deberta-ca に作成しましょう。図 132 では、入出力幅 128 トークン・入出力ベクトル 256 次元・深さ 12 層・アテンションヘッド 4 個・中間ベクトル 768 次元という、やや小さめのモデルにしています。

最後に、my-dir/deberta-ca と UD_Catalan-Ancora から、品詞付与・係り受け解析モデル my-dir/deberta-ca-upos を製作しましょう (図 133)。うまく my-dir/deberta-ca-upos が製作できたら、esupar でテストしてみましょう (図 134)。esupar は縮合語を UPOS 付与においてサポート (図 135) しており、係り受け解析では「al」を 2 語として処理します。

^[123]https://github.com/UniversalDependencies/UD_Catalan-Ancora

^[124]<https://github.com/Softcatala/ca-text-corpus>

```

!pip install esupar spacy-alignments accelerate
import os
url="https://github.com/UniversalDependencies/UD_Catalan-AnCora"
d=os.path.basename(url)
!test -d {d} || git clone --depth=1 {url}
!for F in train dev test ; do cat {d}/*-$$F.conllu > $$F.conllu ; done
url="https://github.com/Softcatala/ca-text-corpus"
d=os.path.basename(url)
!test -d {d} || git clone --depth=1 {url}
!sed -n "s/^# text = //p" *.conllu | cat {d}/data/*.txt - > train.txt

```

図 130: カタルーニャ語ミニモデル製作のための準備

```

from transformers import BertTokenizerFast
from tokenizers import BertWordPieceTokenizer
s=["[CLS]","[PAD]","[SEP]","[UNK]","[MASK]"]
wpt=BertWordPieceTokenizer(strip_accents=False)
wpt.train(files=["train.txt"],vocab_size=30000,special_tokens=s)
wpt.save_model(".")
tkz=BertTokenizerFast(vocab_file="vocab.txt",strip_accents=False,never_split=s)
tkz.save_pretrained("my-dir/tokenizer-ca")

```

図 131: カタルーニャ語向けトークナイザの作成

```

from transformers import (AutoTokenizer,DebertaV2Config,DebertaV2ForMaskedLM,
    DataCollatorForLanguageModeling,TrainingArguments,Trainer)
tkz=AutoTokenizer.from_pretrained("my-dir/tokenizer-ca",model_max_length=128)
t=tkz.convert_tokens_to_ids(["[CLS]","[PAD]","[SEP]","[UNK]","[MASK]"])
cfg=DebertaV2Config(hidden_size=256,num_hidden_layers=12,num_attention_heads=4,
    intermediate_size=768,relative_attention=True,position_biased_input=False,
    pos_att_type=["p2c","c2p"],max_position_embeddings=tkz.model_max_length,
    vocab_size=len(tkz),tokenizer_class=type(tkz).__name__,
    bos_token_id=t[0],pad_token_id=t[1],eos_token_id=t[2])
arg=TrainingArguments(report_to="none",output_dir="/tmp",save_total_limit=2,
    num_train_epochs=3,per_device_train_batch_size=64)
class ReadLineDS(object):
    def __init__(self,file,tokenizer):
        self.tokenizer=tokenizer
        with open(file,"r",encoding="utf-8") as r:
            self.lines=[s.strip() for s in r if s.strip()!=""]
    __len__=lambda self:len(self.lines)
    __getitem__=lambda self,i:self.tokenizer(self.lines[i],truncation=True,
        add_special_tokens=True,max_length=self.tokenizer.model_max_length-2)
trn=Trainer(args=arg,data_collator=DataCollatorForLanguageModeling(tkz),
    model=DebertaV2ForMaskedLM(cfg),train_dataset=ReadLineDS("train.txt",tkz))
trn.train()
trn.save_model("my-dir/deberta-ca")
tkz.save_pretrained("my-dir/deberta-ca")

```

図 132: カタルーニャ語ミニ DeBERTa モデルの作成

```
!python -m esupar.train my-dir/deberta-ca my-dir/deberta-ca-upos .
```

図 133: カタルーニャ語向け品詞付与・係り受け解析ミニモデルの製作

```
!pip install esupar
import esupar
nlp=esupar.load("my-dir/deberta-ca-upos")
doc=nlp("Ho hauries d'aclarir al català i al castellà")
import deplacy
deplacy.serve(doc,port=None)
```

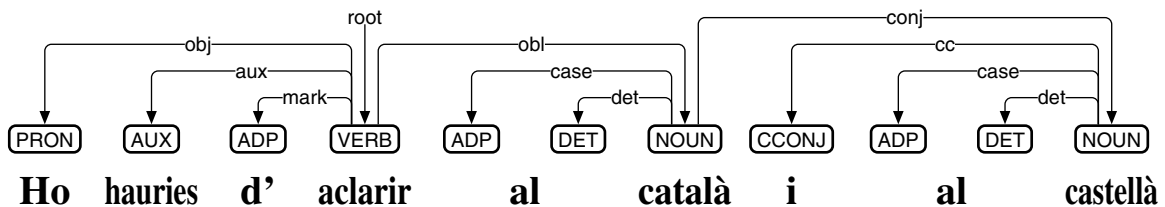


図 134: カタルーニャ語向け品詞付与・係り受け解析ミニモデルのテスト

```
!pip install deplacy transformers
class SeqL(object):
    def __init__(self,bert):
        from transformers import pipeline
        self.tagger=pipeline(task="ner",model=bert)
    def __call__(self,text):
        w=[(t["start"],t["end"],t["entity"],t["word"]) for t in self.tagger(text)]
        u,z="# text = "+text.replace("\n"," ")+"\n","_"
        for i,(s,e,p,t) in enumerate(w,1):
            m,q=z if i<len(w) and e<w[i][0] else "SpaceAfter=No",p.split("|")
            f=z if p.find("=")<0 else "|".join(t for t in q if t.find("=")>0)
            u+="\t".join([str(i),t.strip(),z,q[0].replace(".",","),z,f,z,z,z,m])+"\n"
        return u+"\n"
nlp=SeqL("my-dir/deberta-ca-upos")
doc=nlp("Ho hauries d'aclarir al català i al castellà")
import deplacy
deplacy.serve(doc,port=None)
```



図 135: カタルーニャ語向け品詞付与・係り受け解析ミニモデルでの UPOS 付与

18 ガリシア語編

18.1 ガリシア語 UD における品詞付与と係り受け解析

UDPipe 2 WebAPI は、UDPipe 2 の係り受け解析モデル^[67]を、WebAPI の形で公開しています。内部的には、隣接行列型係り受け解析アルゴリズムを使っています。Google Colaboratory から UDPipe 2 WebAPI の galician-treegal にアクセスするプログラムと、ガリシア語の例文「Deberías acláralo no galego」を品詞付与・係り受け解析した結果を、図 136 に示します。解析結果を見ていくことにしましょう。

```
!pip install deplacy
class UDPipe2WebAPI(object):
    def __init__(self, lang):
        self.url="https://lindat.mff.cuni.cz/services/udpipe/api/process"
        self.params={"model":lang,"tokenizer":"","tagger":"","parser":""}
    def __call__(self, text):
        import urllib.parse, urllib.request, json
        self.params["data"]=urllib.parse.quote(text)
        u=self.url+"?"+"&".join(k+"="+v for k,v in self.params.items())
        with urllib.request.urlopen(u) as r:
            return json.loads(r.read())["result"]
nlp=UDPipe2WebAPI("galician-treegal")
doc=nlp("Deberías acláralo no galego")
import deplacy
deplacy.serve(doc, port=None)
```

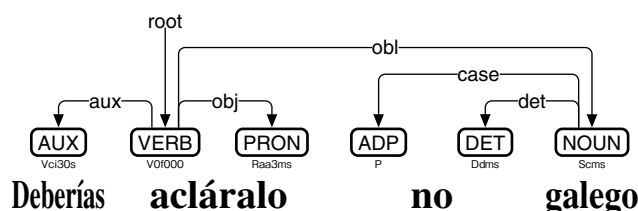


図 136: UDPipe 2 WebAPI によるガリシア語の係り受け解析

「acláralo」は、動詞「aclaraar」と目的語「lo」の2語として扱われていて、objで繋がれています。助動詞「Deberías」はauxで、斜格補語「galego」はoblで、それぞれ繋がれています。「no」は、前置詞「em」と限定詞(冠詞)「o」の2語として扱われていて、それぞれ「galego」からcase・detで繋がられています。

なお、ガリシア語 UD の係り受けタグには、表 4 に加え、受動態を表す nsubj:pass と aux:pass、氏名を表す flat:name、ガリシア語以外を表す flat:foreign も用いられます^[125]。

18.2 ガリシア語ミニ DeBERTa モデルの製作

Google Colaboratory 上の esupar^[47]で、ガリシア語ミニ DeBERTa モデルを製作してみましょう。手順を図 137～140 に示します。

^[125]<https://universaldependencies.org/gl/#syntax>

```
!pip install esupar spacy-alignments accelerate
import os
url="https://github.com/UniversalDependencies/UD_Galician-TreeGal"
!test -d {os.path.basename(url)} || git clone --depth=1 {url}
url="https://github.com/UniversalDependencies/UD_Galician-CTG"
!test -d {os.path.basename(url)} || git clone --depth=1 {url}
!for F in train dev test ; do cat UD_Galician-*/-$$F.conllu > $$F.conllu ; done
!sed -n "s/^# text = //p" *.conllu > train.txt
```

図 137: ガリシア語ミニモデル製作のための準備

```
from transformers import BertTokenizerFast
from tokenizers import BertWordPieceTokenizer
s=[" [CLS] ", " [PAD] ", " [SEP] ", " [UNK] ", " [MASK] "]
wpt=BertWordPieceTokenizer()
wpt.train(files=["train.txt"],vocab_size=3000,special_tokens=s)
wpt.save_model(".")
tkz=BertTokenizerFast(vocab_file="vocab.txt",never_split=s)
tkz.save_pretrained("my-dir/tokenizer-gl")
```

図 138: ガリシア語向けトークナイザの作成

最初に、訓練のための文章を準備しましょう (図 137)。train.txt の文章の材料として、ここでは UD_Galician-TreeGal^[126]と UD_Galician-CTG^[127]を元にしていますが、好みに応じ、ガリシア語の他の文章を増量するのもいいでしょう。なお、各 CoNLL-U ファイルは、品詞付与・係り受け解析の学習にも用います。

続けて、ガリシア語向けトークナイザを、my-dir/tokenizer-gl に作成しましょう。図 138 では、Transformers^[26]の BertTokenizerFast を使って、ガリシア語向けトークナイザを作成しています。語彙数(vocab_size)を 3000 に絞っていますが、もう少し大きくした方がいいかもしれません。

次に train.txt の文章から、DeBERTa モデルを my-dir/deberta-gl に作成しましょう。図 139 では、入出力幅 128 トークン・入出力ベクトル 256 次元・深さ 12 層・アテンションヘッド 4 個・中間ベクトル 768 次元という、やや小さめのモデルにしています。

最後に、my-dir/deberta-gl と UD_Galician-TreeGal・UD_Galician-CTG から、品詞付与・係り受け解析モデル my-dir/deberta-gl-upos を製作しましょう (図 140)。うまく my-dir/deberta-gl-upos が製作できたら、esupar でテストしてみましょう。図 141 の例では「acláralo」が 1 語とみなされており、図 136 とは異なる結果になっているようです。

^[126]https://github.com/UniversalDependencies/UD_Galician-TreeGal

^[127]https://github.com/UniversalDependencies/UD_Galician-CTG

```

from transformers import (AutoTokenizer,DebertaV2Config,DebertaV2ForMaskedLM,
    DataCollatorForLanguageModeling,TrainingArguments,Trainer)
tkz=AutoTokenizer.from_pretrained("my-dir/tokenizer-gl",model_max_length=128)
t=tkz.convert_tokens_to_ids(["[CLS]","[PAD]","[SEP]","[UNK]","[MASK]"])
cfg=DebertaV2Config(hidden_size=256,num_hidden_layers=12,num_attention_heads=4,
    intermediate_size=768,relative_attention=True,position_biased_input=False,
    pos_att_type=["p2c","c2p"],max_position_embeddings=tkz.model_max_length,
    vocab_size=len(tkz),tokenizer_class=type(tkz).__name__,
    bos_token_id=t[0],pad_token_id=t[1],eos_token_id=t[2])
arg=TrainingArguments(report_to="none",output_dir="/tmp",save_total_limit=2,
    num_train_epochs=3,per_device_train_batch_size=64)
class ReadLineDS(object):
    def __init__(self,file,tokenizer):
        self.tokenizer=tokenizer
        with open(file,"r",encoding="utf-8") as r:
            self.lines=[s.strip() for s in r if s.strip()!=""]
        __len__=lambda self:len(self.lines)
        __getitem__=lambda self,i:self.tokenizer(self.lines[i],truncation=True,
            add_special_tokens=True,max_length=self.tokenizer.model_max_length-2)
trn=Trainer(args=arg,data_collator=DataCollatorForLanguageModeling(tkz),
    model=DebertaV2ForMaskedLM(cfg),train_dataset=ReadLineDS("train.txt",tkz))
trn.train()
trn.save_model("my-dir/deberta-gl")
tkz.save_pretrained("my-dir/deberta-gl")

```

図 139: ガリシア語ミニ DeBERTa モデルの作成

```
!python -m esupar.train my-dir/deberta-gl my-dir/deberta-gl-upos .
```

図 140: ガリシア語向け品詞付与・係り受け解析ミニモデルの製作

```

!pip install esupar
import esupar
nlp=esupar.load("my-dir/deberta-gl-upos")
doc=nlp("Deberías aclararlo no galego")
import deplacy
deplacy.serve(doc,port=None)

```

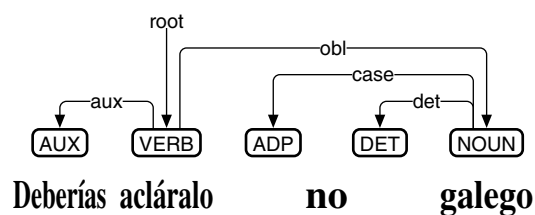


図 141: ガリシア語向け品詞付与・係り受け解析ミニモデルのテスト

19 ポルトガル語編

19.1 ポルトガル語 UD における品詞付与と係り受け解析

modernbert-base-portuguese-ud-embeds^[128]は、隣接行列型係り受け解析アルゴリズムを、事前学習モデルで実装しています。Google Colaboratory 上の Transformers を用いて、ポルトガル語の例文「Você deve esclarecê-lo no português」を modernbert-base-portuguese-ud-embeds で品詞付与・係り受け解析するプログラムと、その結果を図 142 に示します。

```
!pip install deplacy transformers
from transformers import pipeline
mdl="KoichiYasuoka/modernbert-base-portuguese-ud-embeds"
nlp=pipeline("universal-dependencies",mdl,trust_remote_code=True)
doc=nlp("Você deve esclarecê-lo no português")
import deplacy
deplacy.serve(doc,port=None)
```

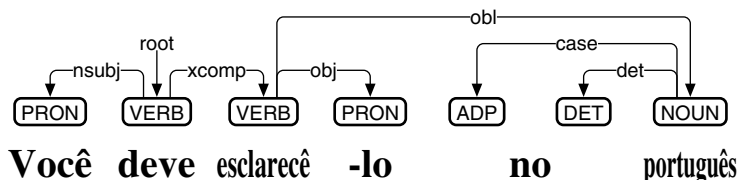


図 142: modernbert-base-portuguese-ud-embeds によるポルトガル語の係り受け解析

「esclarecê-lo」は、動詞「esclarecer」と目的語「lo」の2語として扱われていて、objで繋がっています。動詞「deve」の主語「Você」は、同時に「esclarecer」の主語でもあることから、「Você」がnsubjで、「esclarecer」がxcompで繋がっています。「português」は「esclarecer」の斜格補語であり、oblで繋がっています。「no」は、前置詞「em」と限定詞(冠詞)「o」の2語として扱われていて、それぞれ「português」からcase・detで繋がっています。

なお、ポルトガル語 UD の係り受けタグには、表 4 に加え、受動態を表す nsubj:pass と csubj:pass、所有に関する連体修飾を表す nmod:poss、限定詞の前に置かれる形容詞を表す det:predet、接続における前置副詞を表す cc:preconj も用いられます^[129]。

19.2 ポルトガル語ミニ DeBERTa モデルの製作

Google Colaboratory 上の esupar^[47]で、ポルトガル語ミニ DeBERTa モデルを製作してみましょう。手順を図 143～146 に示します。

最初に、訓練のための文章を準備しましょう(図 143)。train.txt の文章の材料として、ここでは UD_Portuguese-Bosque^[130]と UD_Portuguese-PetroGold^[131]を元にしていますが、好みに応じ、ポルトガル語の他の文章を増量するのもいいでしょう。なお、各 CoNLL-U ファイルは、品詞付与・係り受け解析の学習にも用います。

^[128]<https://huggingface.co/KoichiYasuoka/modernbert-base-portuguese-ud-embeds>

^[129]<https://universaldependencies.org/pt/#syntax>

^[130]https://github.com/UniversalDependencies/UD_Portuguese-Bosque

^[131]https://github.com/UniversalDependencies/UD_Portuguese-PetroGold

```
!pip install esupar spacy-alignments accelerate
import os
url="https://github.com/UniversalDependencies/UD_Portuguese-Bosque"
!test -d {os.path.basename(url)} || git clone --depth=1 {url}
url="https://github.com/UniversalDependencies/UD_Portuguese-PetroGold"
!test -d {os.path.basename(url)} || git clone --depth=1 {url}
!for F in train dev test ; do cat UD_Port*/*-$F.conllu > $$F.conllu ; done
!sed -n "s/^# text = //p" *.conllu > train.txt
```

図 143: ポルトガル語ミニモデル製作のための準備

```
from transformers import BertTokenizerFast
from tokenizers import BertWordPieceTokenizer
s=[" [CLS] ", " [PAD] ", " [SEP] ", " [UNK] ", " [MASK] "]
wpt=BertWordPieceTokenizer(strip_accents=False)
wpt.train(files=["train.txt"],vocab_size=30000,special_tokens=s)
wpt.save_model(".")
tkz=BertTokenizerFast(vocab_file="vocab.txt",strip_accents=False,never_split=s)
tkz.save_pretrained("my-dir/tokenizer-pt")
```

図 144: ポルトガル語向けトークナイザの作成

続けて、ポルトガル語向けトークナイザを、my-dir/tokenizer-pt に作成しましょう。図 144 では、Transformers^[26]の BertTokenizerFast を使って、ポルトガル語向けトークナイザを作成しています。

次に train.txt の文章から、DeBERTa モデルを my-dir/deberta-pt に作成しましょう。図 145 では、入出力幅 128 トークン・入出力ベクトル 256 次元・深さ 12 層・アテンションヘッド 4 個・中間ベクトル 768 次元という、やや小さめのモデルにしています。

最後に、my-dir/deberta-pt と UD_Portuguese-Bosque・UD_Portuguese-PetroGold から、品詞付与・係り受け解析モデル my-dir/deberta-pt-upos を製作しましょう (図 146)。うまく my-dir/deberta-pt-upos が製作できたら、esupar でテストしてみましょう。図 147 の例では、ハイフンが 1 語とみなされており、その周辺の係り受けや品詞が、図 142 とは異なる結果になっているようです。

```

from transformers import (AutoTokenizer,DebertaV2Config,DebertaV2ForMaskedLM,
    DataCollatorForLanguageModeling,TrainingArguments,Trainer)
tkz=AutoTokenizer.from_pretrained("my-dir/tokenizer-pt",model_max_length=128)
t=tkz.convert_tokens_to_ids(["[CLS]","[PAD]","[SEP]","[UNK]","[MASK]"])
cfg=DebertaV2Config(hidden_size=256,num_hidden_layers=12,num_attention_heads=4,
    intermediate_size=768,relative_attention=True,position_biased_input=False,
    pos_att_type=["p2c","c2p"],max_position_embeddings=tkz.model_max_length,
    vocab_size=len(tkz),tokenizer_class=type(tkz).__name__,
    bos_token_id=t[0],pad_token_id=t[1],eos_token_id=t[2])
arg=TrainingArguments(report_to="none",output_dir="/tmp",save_total_limit=2,
    num_train_epochs=3,per_device_train_batch_size=64)
class ReadLineDS(object):
    def __init__(self,file,tokenizer):
        self.tokenizer=tokenizer
        with open(file,"r",encoding="utf-8") as r:
            self.lines=[s.strip() for s in r if s.strip()!=""]
        __len__=lambda self:len(self.lines)
        __getitem__=lambda self,i:self.tokenizer(self.lines[i],truncation=True,
            add_special_tokens=True,max_length=self.tokenizer.model_max_length-2)
trn=Trainer(args=arg,data_collator=DataCollatorForLanguageModeling(tkz),
    model=DebertaV2ForMaskedLM(cfg),train_dataset=ReadLineDS("train.txt",tkz))
trn.train()
trn.save_model("my-dir/deberta-pt")
tkz.save_pretrained("my-dir/deberta-pt")

```

図 145: ポルトガル語ミニ DeBERTa モデルの作成

```
!python -m esupar.train my-dir/deberta-pt my-dir/deberta-pt-upos .
```

図 146: ポルトガル語向け品詞付与・係り受け解析ミニモデルの製作

```

!pip install esupar
import esupar
nlp=esupar.load("my-dir/deberta-pt-upos")
doc=nlp("Você deve esclarecê-lo no português")
import deplacy
deplacy.serve(doc,port=None)

```

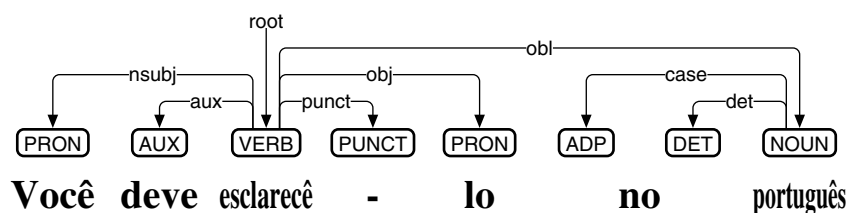


図 147: ポルトガル語向け品詞付与・係り受け解析ミニモデルのテスト

20 イタリア語編

20.1 イタリア語 UD における品詞付与と係り受け解析

UDPipe 2 WebAPI は、UDPipe 2 の係り受け解析モデル^[67]を、WebAPI の形で公開しています。内部的には、隣接行列型係り受け解析アルゴリズムを使っています。Google Colaboratory から UDPipe 2 WebAPI の `italian-isdt` にアクセスするプログラムと、イタリア語の例文「Dovresti chiarirlo nel documento in italiano o latino」を品詞付与・係り受け解析した結果を、図 148 に示します。解析結果を見ていくことにしましょう。

```
!pip install deplacy
class UDPipe2WebAPI(object):
    def __init__(self, lang):
        self.url="https://lindat.mff.cuni.cz/services/udpipe/api/process"
        self.params={"model":lang,"tokenizer":"","tagger":"","parser":""}
    def __call__(self, text):
        import urllib.parse, urllib.request, json
        self.params["data"]=urllib.parse.quote(text)
        u=self.url+"?"+"&".join(k+"="+v for k,v in self.params.items())
        with urllib.request.urlopen(u) as r:
            return json.loads(r.read())["result"]
nlp=UDPipe2WebAPI("italian-isdt")
doc=nlp("Dovresti chiarirlo nel documento in italiano o latino")
import deplacy
deplacy.serve(doc, port=None)
```

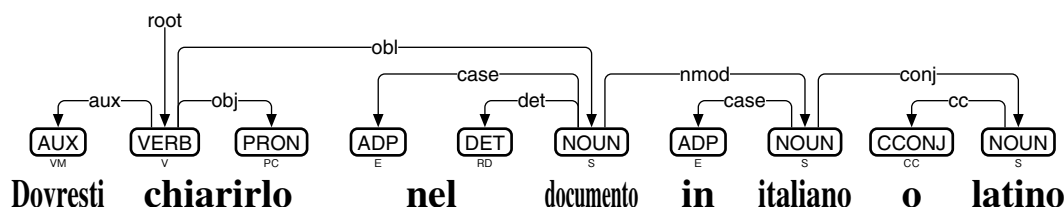


図 148: UDPipe 2 WebAPI によるイタリア語の係り受け解析

「chiarirlo」は、動詞「chiarire」と目的語「lo」の2語として扱われていて、objで繋がっています。助動詞「Dovresti」はauxで、斜格補語「documento」はoblで、それぞれ繋がっています。「nel」は、前置詞「in」と限定詞(冠詞)「il」の2語として扱われていて、それぞれ「documento」からcase・detで繋がっています。「italiano」と「latino」は、いずれも「documento」を修飾していますが、「italiano」の方がnmodで繋がっていて、その先に「latino」がconjで繋がれており、そのまた先にccで「o」が繋がれていると同時に、「italiano」には「in」がcaseで繋がれています。

なお、イタリア語 UD の係り受けタグには、表 4 に加え、受動態を表す `nsubj:pass`・`csubj:pass`・`expl:pass`・`aux:pass`、受動態の動作主を表す `obl:agent`、非人称形式主語を表す `expl:impers`、所有を表す `det:poss`、限定詞の前に置かれる形容詞を表す `det:predet`、関係代名詞による連体修飾節を表す `acl:relcl`、氏名を表す `flat:name`、イタリア語以外を表す `flat:foreign` など、多数のタグが用いられます^[132]。

^[132]<https://universaldependencies.org/it/#syntax>

20.2 イタリア語ミニ DeBERTa モデルの製作

Google Colaboratory 上の esupar^[47]で、イタリア語ミニ DeBERTa モデルを製作してみましょう。手順を図 149～152 に示します。

最初に、訓練のための文章を準備しましょう (図 149)。train.txt の文章の材料として、ここでは UD_Italian-ISDT^[133]と UD_Italian-VIT^[134]を元にしていますが、好みに応じ、イタリア語の他の文章を増量するのもいいでしょう。なお、各 CoNLL-U ファイルは、品詞付与・係り受け解析の学習にも用います。

```
!pip install esupar spacy-alignments accelerate
import os
url="https://github.com/UniversalDependencies/UD_Italian-ISDT"
!test -d {os.path.basename(url)} || git clone --depth=1 {url}
url="https://github.com/UniversalDependencies/UD_Italian-VIT"
!test -d {os.path.basename(url)} || git clone --depth=1 {url}
!for F in train dev test ; do cat UD_Italian-*/-$$F.conllu > $$F.conllu ; done
!sed -n "s/^# text = //p" *.conllu > train.txt
```

図 149: イタリア語ミニモデル製作のための準備

```
from transformers import BertTokenizerFast
from tokenizers import BertWordPieceTokenizer
s=["[CLS]","[PAD]","[SEP]","[UNK]","[MASK]"]
wpt=BertWordPieceTokenizer(strip_accents=False)
wpt.train(files=["train.txt"],vocab_size=30000,special_tokens=s)
wpt.save_model(".")
tkz=BertTokenizerFast(vocab_file="vocab.txt",strip_accents=False,never_split=s)
tkz.save_pretrained("my-dir/tokenizer-it")
```

図 150: イタリア語向けトークナイザの作成

続けて、イタリア語向けトークナイザを、my-dir/tokenizer-it に作成しましょう。図 150 では、Transformers^[26]の BertTokenizerFast を使って、イタリア語向けトークナイザを作成しています。

次に train.txt の文章から、DeBERTa モデルを my-dir/deberta-it に作成しましょう。図 151 では、入出力幅 128 トークン・入出力ベクトル 256 次元・深さ 12 層・アテンションヘッド 4 個・中間ベクトル 768 次元という、やや小さめのモデルにしています。

最後に、my-dir/deberta-it と UD_Italian-ISDT・UD_Italian-VIT から、品詞付与・係り受け解析モデル my-dir/deberta-it-upos を製作しましょう (図 152)。うまく my-dir/deberta-it-upos が製作できたら、esupar でテストしてみましょう。図 153 の例では「chiarirlo」が 1 語とみなされており、図 148 とは異なる結果になっているようです。

^[133]https://github.com/UniversalDependencies/UD_Italian-ISDT

^[134]https://github.com/UniversalDependencies/UD_Italian-VIT

```

from transformers import (AutoTokenizer,DebertaV2Config,DebertaV2ForMaskedLM,
    DataCollatorForLanguageModeling,TrainingArguments,Trainer)
tkz=AutoTokenizer.from_pretrained("my-dir/tokenizer-it",model_max_length=128)
t=tkz.convert_tokens_to_ids(["[CLS]","[PAD]","[SEP]","[UNK]","[MASK]"])
cfg=DebertaV2Config(hidden_size=256,num_hidden_layers=12,num_attention_heads=4,
    intermediate_size=768,relative_attention=True,position_biased_input=False,
    pos_att_type=["p2c","c2p"],max_position_embeddings=tkz.model_max_length,
    vocab_size=len(tkz),tokenizer_class=type(tkz).__name__,
    bos_token_id=t[0],pad_token_id=t[1],eos_token_id=t[2])
arg=TrainingArguments(report_to="none",output_dir="/tmp",save_total_limit=2,
    num_train_epochs=3,per_device_train_batch_size=64)
class ReadLineDS(object):
    def __init__(self,file,tokenizer):
        self.tokenizer=tokenizer
        with open(file,"r",encoding="utf-8") as r:
            self.lines=[s.strip() for s in r if s.strip()!=""]
        __len__=lambda self:len(self.lines)
        __getitem__=lambda self,i:self.tokenizer(self.lines[i],truncation=True,
            add_special_tokens=True,max_length=self.tokenizer.model_max_length-2)
trn=Trainer(args=arg,data_collator=DataCollatorForLanguageModeling(tkz),
    model=DebertaV2ForMaskedLM(cfg),train_dataset=ReadLineDS("train.txt",tkz))
trn.train()
trn.save_model("my-dir/deberta-it")
tkz.save_pretrained("my-dir/deberta-it")

```

図 151: イタリア語ミニ DeBERTa モデルの作成

```
!python -m esupar.train my-dir/deberta-it my-dir/deberta-it-upos .
```

図 152: イタリア語向け品詞付与・係り受け解析ミニモデルの製作

```

!pip install esupar
import esupar
nlp=esupar.load("my-dir/deberta-it-upos")
doc=nlp("Dovresti chiarirlo nel documento in italiano o latino")
import deplacy
deplacy.serve(doc,port=None)

```

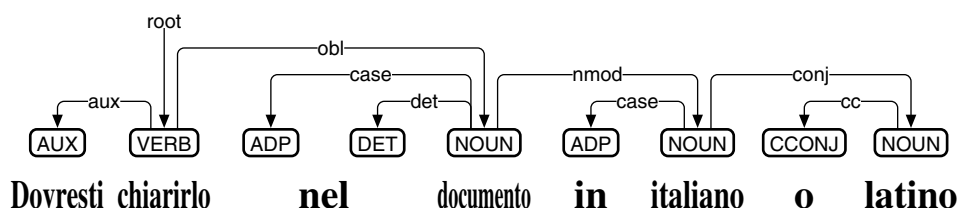


図 153: イタリア語向け品詞付与・係り受け解析ミニモデルのテスト

21 ラテン語編

21.1 ラテン語 UD における品詞付与と係り受け解析

Trankit^[64]は、XLM-RoBERTa をベースに、隣接行列型係り受け解析アルゴリズムを、様々な言語向けに実装しています。Google Colaboratory 上の Trankit を用いて、ラテン語の例文「hoc in documento latino italicove declarare debes」を品詞付与・係り受け解析するプログラムと、その結果を図 154 に示します。解析結果を見ていくことにしましょう。

```
!pip install git+https://github.com/KoichiYasuoka/trankit deplacy
import trankit
nlp=trankit.Pipeline("latin")
doc=nlp("hoc in documento latino italicove declarare debes")
import deplacy
deplacy.serve(doc,port=None)
```

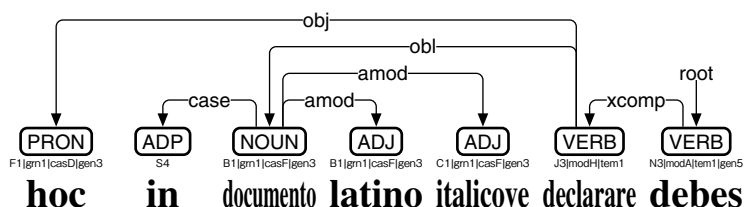


図 154: Trankit によるラテン語の係り受け解析

動詞「debes」の主語(二人称単数)は、動詞「declarare」の意味上の主語でもあることから、これらは **xcomp** で繋がれています。目的語「hoc」は **obj** で、斜格補語「documento」は **obl** で、それぞれ「declarare」に繋がれています。「latino」と「italicove」は、いずれも「documento」を修飾していて **amod** で繋がれており、「in」は **case** で繋がれています。「italicove」を「italico」と「ve」の2語とすべきかどうかは、議論の分かれるところで、ここでは1語とみなしています。

なお、ラテン語 UD の係り受けタグには、表 4 に加え、数多くの拡張が用いられている^[135]のですが、拡張どうしの中で混乱が起こっており、事態の收拾が望まれます。

21.2 ラテン語ミニ DeBERTa モデルの製作

Google Colaboratory 上の esupar^[47]で、ラテン語ミニ DeBERTa モデルを製作してみましょう。手順を図 155～158 に示します。

最初に、訓練のための文章を準備しましょう(図 155)。train.txt の文章の材料として、ここでは UD_Latin-ITTB^[136]と UD_Latin-LLCT^[137]を元にしてはいますが、好みに応じ、ラテン語の他の文章を増量するのもいいでしょう。なお、各 CoNLL-U ファイルは、品詞付与・係り受け解析の学習にも用います。

^[135]<https://universaldependencies.org/la/#relations-overview>

^[136]https://github.com/UniversalDependencies/UD_Latin-ITTB

^[137]https://github.com/UniversalDependencies/UD_Latin-LLCT

```
!pip install esupar spacy-alignments accelerate
import os
url="https://github.com/UniversalDependencies/UD_Latin-ITTB"
!test -d {os.path.basename(url)} || git clone --depth=1 {url}
url="https://github.com/UniversalDependencies/UD_Latin-LLCT"
!test -d {os.path.basename(url)} || git clone --depth=1 {url}
!for F in train dev test ; do cat UD_Latin-*/*-$$F.conllu > $$F.conllu ; done
!sed -n "s/^# text = //p" *.conllu > train.txt
```

図 155: ラテン語ミニモデル製作のための準備

続けて、ラテン語向けトークナイザを、my-dir/tokenizer-la に作成しましょう。図 156 では、Transformers^[26]の BertTokenizerFast を使って、ラテン語向けトークナイザを作成しています。語彙数(vocab_size)を 8000 に絞っていますが、もう少し大きくした方がいいかもしれません。

```
from transformers import BertTokenizerFast
from tokenizers import BertWordPieceTokenizer
s=["[CLS]","[PAD]","[SEP]","[UNK]","[MASK]"]
wpt=BertWordPieceTokenizer(strip_accents=False)
wpt.train(files=["train.txt"],vocab_size=8000,special_tokens=s)
wpt.save_model(".")
tkz=BertTokenizerFast(vocab_file="vocab.txt",strip_accents=False,never_split=s)
tkz.save_pretrained("my-dir/tokenizer-la")
```

図 156: ラテン語向けトークナイザの作成

次に train.txt の文章から、DeBERTa モデルを my-dir/deberta-la に作成しましょう。図 157 では、入出力幅 128 トークン・入出力ベクトル 256 次元・深さ 12 層・アテンションヘッド 4 個・中間ベクトル 768 次元という、やや小さめのモデルにしています。

最後に、my-dir/deberta-la と UD_Latin-ITTB・UD_Latin-LLCT から、品詞付与・係り受け解析モデル my-dir/deberta-la-upos を製作しましょう(図 158)。うまく my-dir/deberta-la-upos が製作できたら、esupar でテストしてみましょう(図 159)。

```

from transformers import (AutoTokenizer,DebertaV2Config,DebertaV2ForMaskedLM,
    DataCollatorForLanguageModeling,TrainingArguments,Trainer)
tkz=AutoTokenizer.from_pretrained("my-dir/tokenizer-la",model_max_length=128)
t=tkz.convert_tokens_to_ids(["[CLS]","[PAD]","[SEP]","[UNK]","[MASK]"])
cfg=DebertaV2Config(hidden_size=256,num_hidden_layers=12,num_attention_heads=4,
    intermediate_size=768,relative_attention=True,position_biased_input=False,
    pos_att_type=["p2c","c2p"],max_position_embeddings=tkz.model_max_length,
    vocab_size=len(tkz),tokenizer_class=type(tkz).__name__,
    bos_token_id=t[0],pad_token_id=t[1],eos_token_id=t[2])
arg=TrainingArguments(report_to="none",output_dir="/tmp",save_total_limit=2,
    num_train_epochs=3,per_device_train_batch_size=64)
class ReadLineDS(object):
    def __init__(self,file,tokenizer):
        self.tokenizer=tokenizer
        with open(file,"r",encoding="utf-8") as r:
            self.lines=[s.strip() for s in r if s.strip()!=""]
        __len__=lambda self:len(self.lines)
        __getitem__=lambda self,i:self.tokenizer(self.lines[i],truncation=True,
            add_special_tokens=True,max_length=self.tokenizer.model_max_length-2)
trn=Trainer(args=arg,data_collator=DataCollatorForLanguageModeling(tkz),
    model=DebertaV2ForMaskedLM(cfg),train_dataset=ReadLineDS("train.txt",tkz))
trn.train()
trn.save_model("my-dir/deberta-la")
tkz.save_pretrained("my-dir/deberta-la")

```

図 157: ラテン語ミニ DeBERTa モデルの作成

```
!python -m esupar.train my-dir/deberta-la my-dir/deberta-la-upos .
```

図 158: ラテン語向け品詞付与・係り受け解析ミニモデルの製作

```

!pip install esupar
import esupar
nlp=esupar.load("my-dir/deberta-la-upos")
doc=nlp("hoc in documento latino italicove declarare debes")
import deplacy
deplacy.serve(doc,port=None)

```

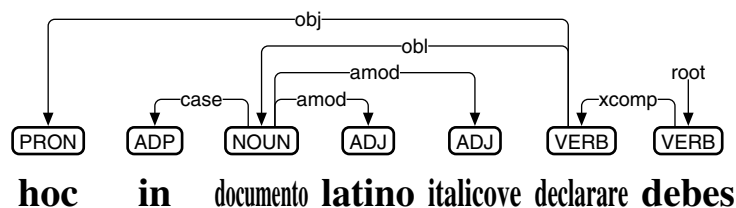


図 159: ラテン語向け品詞付与・係り受け解析ミニモデルのテスト

22 バスク語編

22.1 バスク語 UD における品詞付与と係り受け解析

バスク語は、絶対格 (absolutibo) と能格 (ergatibo) と与格 (datibo) で動詞の項を示す能格言語です。動詞の項が1つの場合は絶対格を、2つの場合は絶対格と能格を、3つの場合は絶対格と能格と与格を使います。ただし、バスク語 UD においては、動詞の項が1つの場合は絶対格を主語とみなし、動詞の項が2つ以上の場合は絶対格を目的語とみなす (能格を主語とみなし、与格を間接目的語とみなす)^{[138][139]} 点で、むしろラテン語 UD やスペイン語 UD との連絡性を重視しています。

Trankit^[64] は、XLM-RoBERTa をベースに、隣接行列型係り受け解析アルゴリズムを、様々な言語向けに実装しています。Google Colaboratory 上の Trankit を用いて、バスク語の例文「Gaur ama hil zen gaur goizean ama hil nuen」を品詞付与・係り受け解析するプログラムと、その結果を図 160 に示します。解析結果を見ていくことにしましょう。

```
!pip install git+https://github.com/KoichiYasuoka/trankit deplacy
import trankit
nlp=trankit.Pipeline("basque")
doc=nlp("Gaur ama hil zen gaur goizean ama hil nuen")
import deplacy
deplacy.serve(doc,port=None)
```

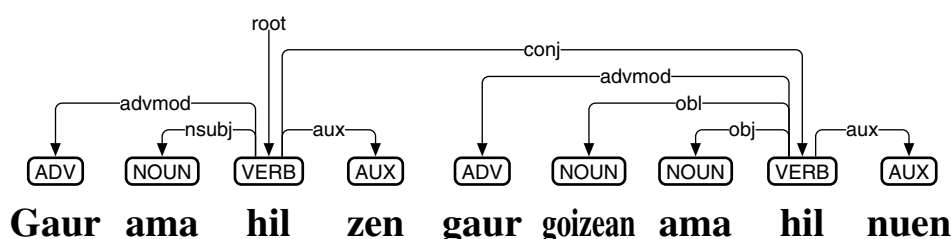


図 160: Trankit によるバスク語の係り受け解析

2つの動詞「hil」は conj で繋がれており、それぞれに絶対格の名詞「ama」が繋がっています。前者は「zen」(izan の3人称単数・直接法過去)が aux で繋がれていることから、絶対格「ama」は主語 (nsubj) とみなされます。後者は「nuen」(izan の絶対格3人称単数・能格1人称単数・直接法過去)が aux で繋がれていることから、能格(1人称単数)は「nuen」に内包されているとみなされ、絶対格「ama」は目的語 (obj) とみなされます。名詞「goizean」(goiz の内格 (inesibo) 単数) は、斜格補語 (obl) です。副詞「gaur」は、それぞれ advmod で繋がれています。

^[138]Izaskun Aldezabal, Maria Jesus Aranzabe, Jose Mari Arriola, and Arantza Diaz de Ilarraza: Syntactic annotation in the Reference Corpus for the Processing of Basque (EPEC): Theoretical and practical issues, Corpus Linguistics and Linguistic Theory, Vol.5, No.2 (October 2009), pp.241-269.

^[139]Maria Jesus Aranzabe, Aitziber Atutxa, Kepa Bengoetxea, Arantza Diaz de Ilarraza, Iakes Goenaga, Koldo Gojenola, Larraitz Uri: Dependenzia Unibertsalen eredura egokitutako euskarazko zuhaitz-bankua, Ekaia, zk:35 (2019), pp.291-307.

22.2 バスク語ミニ DeBERTa モデルの製作

Google Colaboratory 上の esupar^[47]で、バスク語ミニ DeBERTa モデルを製作してみましょう。手順を図 161～164 に示します。

最初に、訓練のための文章を準備しましょう (図 161)。train.txt の文章の材料として、ここでは UD_Basque-BDT^[140]と BasqueParl コーパス^[141]を元にしていますが、好みに応じ、バスク語の他の文章を増量するのもいいでしょう。なお、UD_Basque-BDT の各 CoNLL-U ファイルは、品詞付与・係り受け解析の学習にも用います。

```
!pip install esupar accelerate datasets
import os,datasets
url="https://github.com/UniversalDependencies/UD_Basque-BDT"
d=os.path.basename(url)
!test -d {d} || git clone --depth=1 {url}
!for F in train dev test ; do cat {d}/*-$$F.conllu > $$F.conllu ; done
c=[k["text"] for k in datasets.load_dataset("HiTZ/basqueparl")["train"]]
with open("train.txt","w",encoding="utf-8") as w:
    print("\n".join(c),file=w)
!sed -n "s/^# text = //p" *.conllu >> train.txt
```

図 161: バスク語ミニモデル製作のための準備

```
from transformers import BertTokenizerFast
from tokenizers import BertWordPieceTokenizer
s=["[CLS]","[PAD]","[SEP]","[UNK]","[MASK]"]
wpt=BertWordPieceTokenizer(strip_accents=False)
wpt.train(files=["train.txt"],vocab_size=30000,special_tokens=s)
wpt.save_model(".")
tkz=BertTokenizerFast(vocab_file="vocab.txt",strip_accents=False,never_split=s)
tkz.save_pretrained("my-dir/tokenizer-eu")
```

図 162: バスク語向けトークナイザの作成

続けて、バスク語向けトークナイザを、my-dir/tokenizer-eu に作成しましょう。図 162 では、Transformers^[26]の BertTokenizerFast を使って、バスク語向けトークナイザを作成しています。

次に train.txt の文章から、DeBERTa モデルを my-dir/deberta-eu に作成しましょう。図 163 では、入出力幅 128 トークン・入出力ベクトル 256 次元・深さ 12 層・アテンションヘッド 4 個・中間ベクトル 768 次元という、やや小さめのモデルにしています。

最後に、my-dir/deberta-eu と UD_Basque-BDT から、品詞付与・係り受け解析モデル my-dir/deberta-eu-upos を製作しましょう (図 164)。うまく my-dir/deberta-eu-upos が製作できたら、esupar でテストしてみましょう (図 165)。

^[140]https://github.com/UniversalDependencies/UD_Basque-BDT

^[141]Nayla Escribano, Jon Ander González, Julen Orbegoza-Terradillos, Ainara Larrondo-Ureta, Simón Peña-Fernández, Olatz Perez-de-Viñaspre and Rodrigo Agerri: BasqueParl: A Bilingual Corpus of Basque Parliamentary Transcriptions, LREC 2022: Thirteenth Conference on Language Resources and Evaluation (June 2022), pp.3382-3390.

```

from transformers import (AutoTokenizer,DebertaV2Config,DebertaV2ForMaskedLM,
    DataCollatorForLanguageModeling,TrainingArguments,Trainer)
tkz=AutoTokenizer.from_pretrained("my-dir/tokenizer-eu",model_max_length=128)
t=tkz.convert_tokens_to_ids(["[CLS]","[PAD]","[SEP]","[UNK]","[MASK]"])
cfg=DebertaV2Config(hidden_size=256,num_hidden_layers=12,num_attention_heads=4,
    intermediate_size=768,relative_attention=True,position_biased_input=False,
    pos_att_type=["p2c","c2p"],max_position_embeddings=tkz.model_max_length,
    vocab_size=len(tkz),tokenizer_class=type(tkz).__name__,
    bos_token_id=t[0],pad_token_id=t[1],eos_token_id=t[2])
arg=TrainingArguments(report_to="none",output_dir="/tmp",save_total_limit=2,
    num_train_epochs=3,per_device_train_batch_size=64)
class ReadLineDS(object):
    def __init__(self,file,tokenizer):
        self.tokenizer=tokenizer
        with open(file,"r",encoding="utf-8") as r:
            self.lines=[s.strip() for s in r if s.strip()!=""]
        __len__=lambda self:len(self.lines)
        __getitem__=lambda self,i:self.tokenizer(self.lines[i],truncation=True,
            add_special_tokens=True,max_length=self.tokenizer.model_max_length-2)
trn=Trainer(args=arg,data_collator=DataCollatorForLanguageModeling(tkz),
    model=DebertaV2ForMaskedLM(cfg),train_dataset=ReadLineDS("train.txt",tkz))
trn.train()
trn.save_model("my-dir/deberta-eu")
tkz.save_pretrained("my-dir/deberta-eu")

```

図 163: バスク語ミニ DeBERTa モデルの作成

```
!python -m esupar.train my-dir/deberta-eu my-dir/deberta-eu-upos .
```

図 164: バスク語向け品詞付与・係り受け解析ミニモデルの製作

```

!pip install esupar
import esupar
nlp=esupar.load("my-dir/deberta-eu-upos")
doc=nlp("Gaur ama hil zen gaur goizean ama hil nuen")
import deplacy
deplacy.serve(doc,port=None)

```

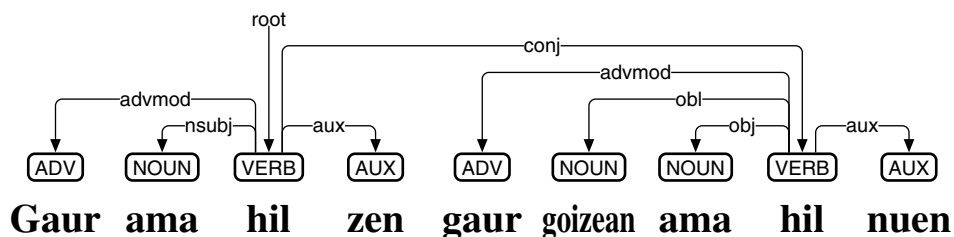


図 165: バスク語向け品詞付与・係り受け解析ミニモデルのテスト

23 トルコ語編

23.1 トルコ語 UD における品詞付与と係り受け解析

UDPipe 2 WebAPI は、UDPipe 2 の係り受け解析モデル^[67]を、WebAPI の形で公開しています。内部的には、隣接行列型係り受け解析アルゴリズムを使っています。Google Colaboratory から UDPipe 2 WebAPI の `turkish-kenet` にアクセスするプログラムと、トルコ語の例文「Ay dağın diğer tarafında yükseldi」を品詞付与・係り受け解析した結果を、図 166 に示します。解析結果を見ていくことにしましょう。

```
!pip install deplacy
class UDPipe2WebAPI(object):
    def __init__(self, lang):
        self.url="https://lindat.mff.cuni.cz/services/udpipe/api/process"
        self.params={"model":lang,"tokenizer":"","tagger":"","parser":""}
    def __call__(self, text):
        import urllib.parse, urllib.request, json
        self.params["data"]=urllib.parse.quote(text)
        u=self.url+"?"+"&".join(k+"="+v for k,v in self.params.items())
        with urllib.request.urlopen(u) as r:
            return json.loads(r.read())["result"]
nlp=UDPipe2WebAPI("turkish-kenet")
doc=nlp("Ay dağın diğer tarafında yükseldi")
import deplacy
deplacy.serve(doc, port=None)
```

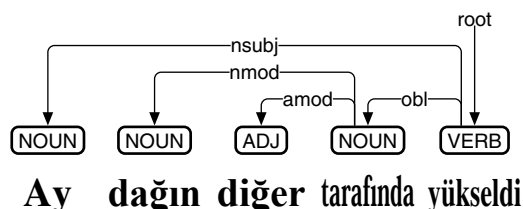


図 166: UDPipe 2 WebAPI によるトルコ語の係り受け解析

動詞「yükseldi」の主語は「Ay」であり、`nsubj`で繋がられています。「tarafında」は、場所を表す接尾辞 (-da) を伴っており、`obl`(斜格補語)で繋がられています。「tarafında」には、「dağın」(dağ の属格)が `nmod` で、「diğer」が `amod` で、それぞれ繋がられています。

なお、トルコ語 UD の係り受けタグには、表 4 に加え、数多くの拡張が用いられている^[142]のですが、拡張に関する情報が不十分で、混乱が発生しているようです。

23.2 トルコ語ミニ DeBERTa モデルの製作

Google Colaboratory 上の `esupar`^[47]で、トルコ語ミニ DeBERTa モデルを製作してみましょう。手順を図 167～170 に示します。

^[142]<https://universaldependencies.org/tr/#relations-overview>

```

!pip install esupar accelerate
import os
url="https://github.com/UniversalDependencies/UD_Turkish-Kenet"
!test -d {os.path.basename(url)} || git clone --depth=1 {url}
url="https://github.com/UniversalDependencies/UD_Turkish-Penn"
!test -d {os.path.basename(url)} || git clone --depth=1 {url}
!for F in train dev test ; do cat UD_Turkish-*/*-$$F.conllu > $$F.conllu ; done
!sed -n "s/^# text = //p" *.conllu > train.txt

```

図 167: トルコ語ミニモデル製作のための準備

最初に、訓練のための文章を準備しましょう (図 167)。train.txt の文章の材料として、ここでは UD_Turkish-Kenet^[143]と UD_Turkish-Penn^[144]を元にしていますが、好みに応じ、トルコ語の他の文章を増量するのもいいでしょう。なお、各 CoNLL-U ファイルは、品詞付与・係り受け解析の学習にも用います。

続けて、トルコ語向けトークナイザを、my-dir/tokenizer-tr に作成しましょう。図 168 では、Transformers^[26]の BertTokenizerFast を使って、トルコ語向けトークナイザを作成しています。語彙数(vocab_size)を 8000 に絞っていますが、もう少し大きくした方がいいかもしれません。

```

from transformers import BertTokenizerFast
from tokenizers import BertWordPieceTokenizer
s=[" [CLS]", " [PAD]", " [SEP]", " [UNK]", " [MASK]"]
wpt=BertWordPieceTokenizer(strip_accents=False)
wpt.train(files=["train.txt"],vocab_size=8000,special_tokens=s)
wpt.save_model(".")
tkz=BertTokenizerFast(vocab_file="vocab.txt",strip_accents=False,never_split=s)
tkz.save_pretrained("my-dir/tokenizer-tr")

```

図 168: トルコ語向けトークナイザの作成

次に train.txt の文章から、DeBERTa モデルを my-dir/deberta-tr に作成しましょう。図 169 では、入出力幅 128 トークン・入出力ベクトル 256 次元・深さ 12 層・アテンションヘッド 4 個・中間ベクトル 768 次元という、やや小さめのモデルにしています。

最後に、my-dir/deberta-tr と UD_Turkish-Kenet・UD_Turkish-Penn から、品詞付与・係り受け解析モデル my-dir/deberta-tr-upos を製作しましょう (図 170)。うまく my-dir/deberta-tr-upos が製作できたら、esupar でテストしてみましょう (図 171)。

^[143]https://github.com/UniversalDependencies/UD_Turkish-Kenet

^[144]https://github.com/UniversalDependencies/UD_Turkish-Penn

```

from transformers import (AutoTokenizer,DebertaV2Config,DebertaV2ForMaskedLM,
    DataCollatorForLanguageModeling,TrainingArguments,Trainer)
tkz=AutoTokenizer.from_pretrained("my-dir/tokenizer-tr",model_max_length=128)
t=tkz.convert_tokens_to_ids(["[CLS]","[PAD]","[SEP]","[UNK]","[MASK]"])
cfg=DebertaV2Config(hidden_size=256,num_hidden_layers=12,num_attention_heads=4,
    intermediate_size=768,relative_attention=True,position_biased_input=False,
    pos_att_type=["p2c","c2p"],max_position_embeddings=tkz.model_max_length,
    vocab_size=len(tkz),tokenizer_class=type(tkz).__name__,
    bos_token_id=t[0],pad_token_id=t[1],eos_token_id=t[2])
arg=TrainingArguments(report_to="none",output_dir="/tmp",save_total_limit=2,
    num_train_epochs=3,per_device_train_batch_size=64)
class ReadLineDS(object):
    def __init__(self,file,tokenizer):
        self.tokenizer=tokenizer
        with open(file,"r",encoding="utf-8") as r:
            self.lines=[s.strip() for s in r if s.strip()!=""]
        __len__=lambda self:len(self.lines)
        __getitem__=lambda self,i:self.tokenizer(self.lines[i],truncation=True,
            add_special_tokens=True,max_length=self.tokenizer.model_max_length-2)
trn=Trainer(args=arg,data_collator=DataCollatorForLanguageModeling(tkz),
    model=DebertaV2ForMaskedLM(cfg),train_dataset=ReadLineDS("train.txt",tkz))
trn.train()
trn.save_model("my-dir/deberta-tr")
tkz.save_pretrained("my-dir/deberta-tr")

```

図 169: トルコ語ミニ DeBERTa モデルの作成

```
!python -m esupar.train my-dir/deberta-tr my-dir/deberta-tr-upos .
```

図 170: トルコ語向け品詞付与・係り受け解析ミニモデルの製作

```

!pip install esupar
import esupar
nlp=esupar.load("my-dir/deberta-tr-upos")
doc=nlp("Ay dağın diğer tarafında yükseldi")
import deplacy
deplacy.serve(doc,port=None)

```

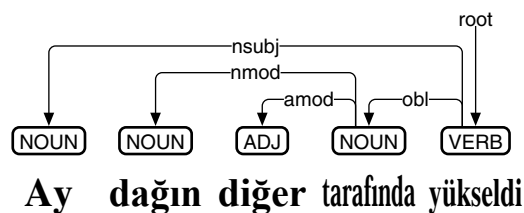


図 171: トルコ語向け品詞付与・係り受け解析ミニモデルのテスト

24 韓国語編

24.1 韓国語 UD における品詞付与と係り受け解析

書写言語としての韓国語は、ハングルが語節 (어절) ごとに分かち書きされています。これにしたがって韓国語 UD では、語節ごとに品詞付与をおこない、語節と語節の係り受けを定義しています。

esupar^[47]は、隣接行列型係り受け解析アルゴリズムを、事前学習モデルで実装しています。Google Colaboratory 上の esupar を用いて、roberta-base-korean-upos で韓国語の例文「한국어는 한자를 사용하지 않고 한글만으로 쓴다」を係り受け解析するプログラムと、その結果を図 172 に示します。解析結果を見ていくことにしましょう。

```
!pip install esupar
import esupar
nlp=esupar.load("KoichiYasuoka/roberta-base-korean-upos")
doc=nlp("한국어는 한자를 사용하지 않고 한글만으로 쓴다")
import deplacy
deplacy.serve(doc,port=None)
```

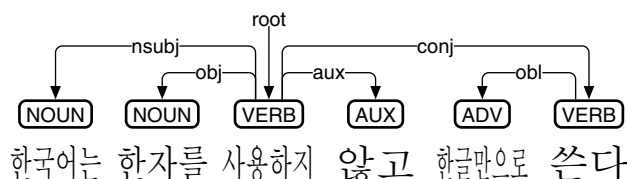


図 172: roberta-base-korean-upos による韓国語の係り受け解析

この例文は、「사용하지」と「쓴다」を中心とする2つの節から成っており、conjで繋がれています。述語「사용하지」からは、主語「한국어는」がnsubjで、目的語「한자를」がobjで、否定語「않고」がauxで、それぞれ繋がられています。述語「쓴다」からは、斜

```
!pip install esupar
import esupar
nlp=esupar.load("KoichiYasuoka/roberta-base-korean-upos")
doc=nlp("한국어는 한자(漢字)를 사용하지 않고 한글만으로 쓴다")
import deplacy
deplacy.serve(doc,port=None)
```

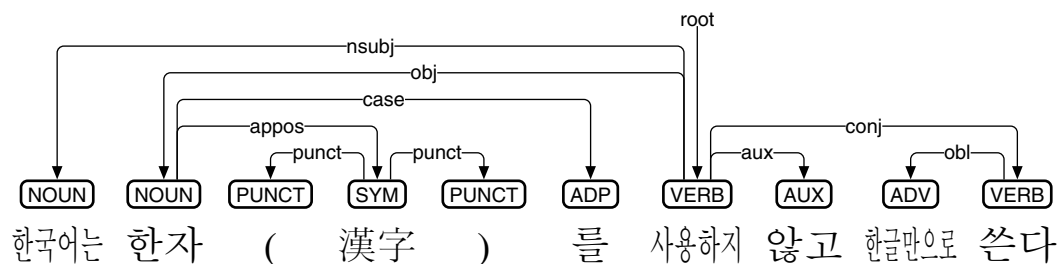


図 173: カッコ書きの漢字を含む韓国語の係り受け解析

韓国語の文には、しばしばカッコ書きの漢字が挿入されます。たとえば「한자를」という語節に「한자(漢字)를」という形で、漢字が挿入されます。この場合、韓国語 UD ではカッコ書きの部分を独立して扱う^[145]ため、結果的に「한자」と「를」は分断されて、形態素としての係り受け解析がおこなわれます(図 173)。つまり、このような文においては、語節と形態素が入り混じった形で係り受け解析をおこなう必要があり、どうしても解析精度が下がってしまうのです。

一方で、語節ではなく形態素のみを用いて、韓国語の係り受け解析をおこなうアプローチ^[146]も可能です。Google Colaboratory 上の esupar を用いて、roberta-base-korean-morphupos で韓国語の例文「한국어는 한자(漢字)를 사용하지 않고 한글만으로 쓴다」を係り受け解析するプログラムと、その結果を図 174 に示します。

한국어는 한자(漢字)를 사용하하지 않음

고 한글만 으로 쓴다

述語を形態素に分解すると、品詞付与が複雑になります。述語「사용하지」の「사용」は名詞なので **NOUN** とすべきですが、その場合に「하」と「지」の UPOS 品詞が **PART**

[146] Yige Chen, Eunkyul Leah Jo, Yundong Yao, KyungTae Lim, Miikka Silfverberg, Francis M. Tyers, Jungyeul Park: Yet Another Format of Universal Dependencies for Korean, Proceedings of the 29th International Conference on Computational Linguistics (October 2022), pp.5432-5437.

と CCONJ でいいのか、議論の余地があります。あるいは述語「쓰다」を形態素に分解すると「쓰」「다」なので、ハングルの途中で形態素の境界が来てしまいます。韓国語の形態素にもとづく品詞付与と係り受け解析は、理論的には可能なものの、述語の処理単位としては複雑すぎて、あまり実用的ではないようです。

24.3 韓国語ミニ DeBERTa モデルの製作

Google Colaboratory 上の esupar で、韓国語ミニ DeBERTa モデルを製作してみましょう。手順を図 175～178 に示します。

最初に、訓練のための文章を準備しましょう (図 175)。train.txt の文章の材料として、ここでは UD_Korean-Kaist^[147]と UD_Korean-GSD^[148]を元にしていますが、好みに応じ、韓国語の他の文章を増量するのもいいでしょう。なお、各 CoNLL-U ファイルは、品詞付与・係り受け解析の学習にも用います。

```
!pip install esupar spacy-alignments accelerate
import os
url="https://github.com/UniversalDependencies/UD_Korean-Kaist"
!test -d {os.path.basename(url)} || git clone --depth=1 {url}
url="https://github.com/UniversalDependencies/UD_Korean-GSD"
!test -d {os.path.basename(url)} || git clone --depth=1 {url}
!for F in train dev test ; do cat UD_Korean-*/-$$F.conllu > $$F.conllu ; done
!sed -n "s/^# text = //p" *.conllu > train.txt
```

図 175: 韓国語ミニモデル製作のための準備

```
from transformers import BertTokenizerFast
from tokenizers import pre_tokenizers, normalizers
from esupar.hangul import hangul
s, t = ["[CLS]", "[PAD]", "[SEP]", "[UNK]", "[MASK]"], {}
for k, v in hangul.items():
    for c in "".join(v).replace("(", "").replace(")", ""):
        t[c] = k
with open("train.txt", "r", encoding="utf-8") as r:
    v = set(t[c] if c in t else c for c in r.read() if not c.isspace())
with open("vocab.txt", "w", encoding="utf-8") as w:
    print("\n".join(s + sorted(v)) + "\n##".join([""] + sorted(v)), file=w)
tkz = BertTokenizerFast(vocab_file="vocab.txt", never_split=s,
    do_lower_case=False, strip_accents=False, tokenize_chinese_chars=False)
tkz.backend_tokenizer.pre_tokenizer = pre_tokenizers.Whitespace()
tkz.backend_tokenizer.normalizer = normalizers.Sequence([normalizers.Nmt()] +
    [normalizers.Replace(k, v) for k, v in t.items()] + [normalizers.NFKC()])
tkz.save_pretrained("my-dir/tokenizer-ko")
```

図 176: 韓国語向け単文字トークナイザの作成

^[147]https://github.com/UniversalDependencies/UD_Korean-Kaist

^[148]https://github.com/UniversalDependencies/UD_Korean-GSD

続けて、韓国語向け単文字トークナイザを、my-dir/tokenizer-ko に作成しましょう。図 176 では、Transformers^[26]の BertTokenizerFast を使って、韓国語向け単文字トークナイザを作成しています。なお、漢字はハングルに変換しています。

次に train.txt の文章から、DeBERTa モデルを my-dir/deberta-ko に作成しましょう。図 177 では、入出力幅 128 トークン・入出力ベクトル 256 次元・深さ 12 層・アテンションヘッド 4 個・中間ベクトル 768 次元という、やや小さめのモデルにしています。

最後に、my-dir/deberta-ko と UD_Korean-Kaist・UD_Korean-GSD から、品詞付与・係り受け解析モデル my-dir/deberta-ko-upos を製作しましょう (図 178)。うまく my-dir/deberta-ko-upos が製作できたら、esupar でテストしてみましょう (図 179)。

```

from transformers import (AutoTokenizer,DebertaV2Config,DebertaV2ForMaskedLM,
    DataCollatorForLanguageModeling,TrainingArguments,Trainer)
tkz=AutoTokenizer.from_pretrained("my-dir/tokenizer-ko",model_max_length=128)
t=tkz.convert_tokens_to_ids(["[CLS]","[PAD]","[SEP]","[UNK]","[MASK]"])
cfg=DebertaV2Config(hidden_size=256,num_hidden_layers=12,num_attention_heads=4,
    intermediate_size=768,relative_attention=True,position_biased_input=False,
    pos_att_type=["p2c","c2p"],max_position_embeddings=tkz.model_max_length,
    vocab_size=len(tkz),tokenizer_class=type(tkz).__name__,
    bos_token_id=t[0],pad_token_id=t[1],eos_token_id=t[2])
arg=TrainingArguments(report_to="none",output_dir="/tmp",save_total_limit=2,
    num_train_epochs=3,per_device_train_batch_size=64)
class ReadLineDS(object):
    def __init__(self,file,tokenizer):
        self.tokenizer=tokenizer
        with open(file,"r",encoding="utf-8") as r:
            self.lines=[s.strip() for s in r if s.strip()!=""]
    __len__=lambda self:len(self.lines)
    __getitem__=lambda self,i:self.tokenizer(self.lines[i],truncation=True,
        add_special_tokens=True,max_length=self.tokenizer.model_max_length-2)
trn=Trainer(args=arg,data_collator=DataCollatorForLanguageModeling(tkz),
    model=DebertaV2ForMaskedLM(cfg),train_dataset=ReadLineDS("train.txt",tkz))
trn.train()
trn.save_model("my-dir/deberta-ko")
tkz.save_pretrained("my-dir/deberta-ko")

```

図 177: 韓国語ミニ DeBERTa モデルの作成

```
!python -m esupar.train my-dir/deberta-ko my-dir/deberta-ko-upos .
```

図 178: 韓国語向け品詞付与・係り受け解析ミニモデルの製作

```

!pip install esupar
import esupar
nlp=esupar.load("my-dir/deberta-ko-upos")
doc=nlp("한국어는 한자(漢字)를 사용하지 않고 한글만으로 쓴다")
import deplacy
deplacy.serve(doc,port=None)

```

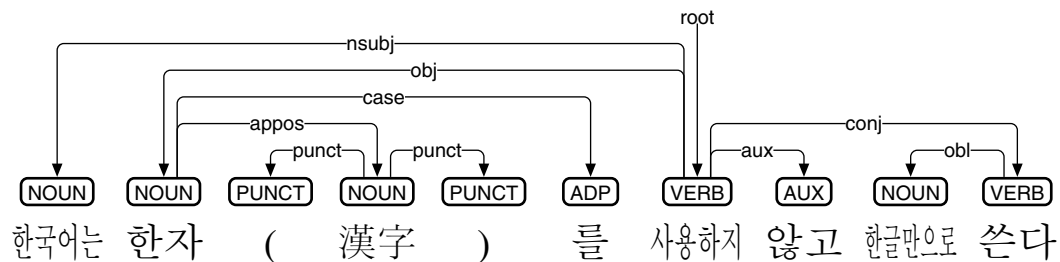


図 179: 韓国語向け品詞付与・係り受け解析ミニモデルのテスト

25 アイヌ語編

25.1 アイヌ語UDにおける品詞付与と係り受け解析

esupar^[47]は、隣接行列型係り受け解析アルゴリズムを、事前学習モデルで実装しています。Google Colaboratory 上の esupar を用いて、deberta-base-ainu-upos でアイヌ語の例文「アイヌイタク アニ エネ アカラ ヒ エチヌカレ」を係り受け解析するプログラムと、その結果を図 180 に示します。解析結果を見ていくことにしましょう。

```
!pip install esupar
import esupar
nlp=esupar.load("KoichiYasuoka/deberta-base-ainu-upos")
doc=nlp("アイヌイタク アニ エネ アカラ ヒ エチヌカレ")
import deplacy
deplacy.serve(doc,port=None)
```

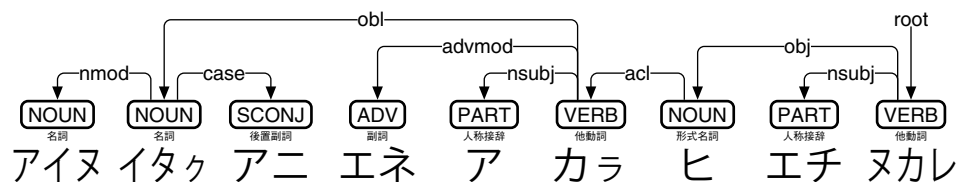


図 180: deberta-base-ainu-upos によるアイヌ語 (カタカナ) の係り受け解析

複他動詞「ヌカレ」は、本来は3つの項(主語・目的語・間接目的語)を取るのですが、ここでは「エチ」が主語(1人称複数)と間接目的語(2人称複数)を兼ねています。したがって「エチ」は nsubj と iobj の両方で繋ぐべきなのですが、UDでは各単語に入るリンクは1つだけなので nsubj を優先しています。目的語「ヒ」は obj で繋がれています。「アカラ」は「ヒ」を体言修飾していて、「カラ」と「ヒ」が acl で繋がれており、「ア」は nsubj で繋がれています。「エネ」は advmod で、複合名詞の「アイヌイタク」は obl で繋がれています。「アイヌ」と「イタク」は nmod で繋がれた上に、「アニ」が case で繋がれています。

```
!pip install esupar
import esupar
nlp=esupar.load("KoichiYasuoka/deberta-base-ainu-upos")
doc=nlp("aynu itak ani ene a=kar hi eci=nukare")
import deplacy
deplacy.serve(doc,port=None)
```

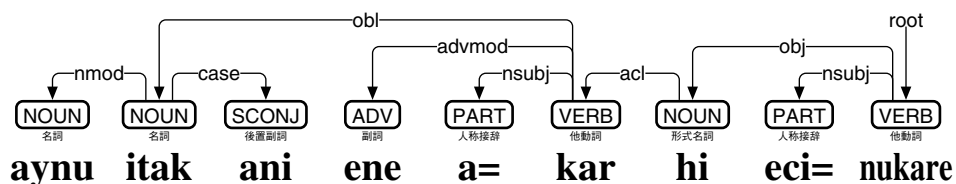


図 181: deberta-base-ainu-upos によるアイヌ語 (ラテン文字) の係り受け解析

```
!pip install esupar
import esupar
nlp=esupar.load("KoichiYasuoka/deberta-base-ainu-upos")
doc=nlp("айну итак ани энэ акар хи эчинукарэ")
import deplacy
deplacy.serve(doc,port=None)
```

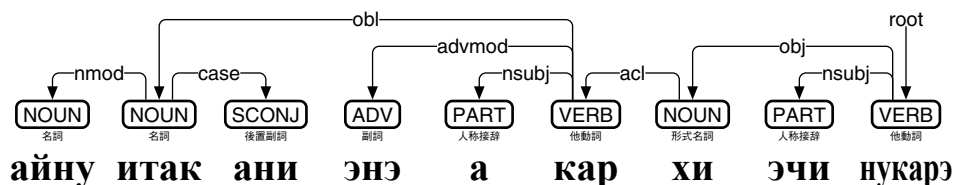


図 182: deberta-base-ainu-upos によるアイヌ語 (キリル文字) の係り受け解析

deberta-base-ainu-upos は、カタカナで書かれたアイヌ語のみならず、ラテン文字やキリル文字で書かれたアイヌ語も解析可能^[149]です。Google Colaboratory 上の esupar を用いて、deberta-base-ainu-upos でアイヌ語の例文「aynu itak ani ene a=kar hi eci=nukare」を係り受け解析するプログラムと、その結果を図 181 に示します。同様に「айну итак ани энэ акар хи эчинукарэ」を係り受け解析するプログラムと、その結果を図 182 に示します。

25.2 アイヌ語ミニ DeBERTa モデルの製作

Google Colaboratory 上の esupar で、アイヌ語ミニ DeBERTa モデルを製作してみましょう。手順を図 183～187 に示します。

最初に、訓練のための文章を準備しましょう (図 183)。train.txt の文章の材料として、ここでは国立国語研究所アイヌ語口承文芸コーパス^[150]と UD-Ainu^[151]を元にはしていますが、

```
!pip install esupar pdfminer accelerate
import os
url="https://ainu.ninjal.ac.jp/static/pdf/text_jp/all.text.jp.pdf"
f=os.path.basename(url)
!test -f {f} || curl -LO {url}
s="y/\t/ /;s/^[0-9/]*//;/^ *$/d;/^[!-?A-Za-z ]*$/p;d"
!pdf2txt.py {f} | LANG=C sed "{s}" > train.txt
url="https://github.com/KoichiYasuoka/UD-Ainu"
d=os.path.basename(url)
!test -d {d} || git clone --depth=1 {url}
t='''BEGIN{FS=OFS="\t"};if(NF==10){u=u$0"\n";if($1~/-/){printf("%s ",$3);
$2=$3;t=t$0"\n"}}else if($0==""){print;print u"\n"t>"train.upos";
if(u~/\t0\troot\t/){print u>(NR%10?"train":"test").conllu;u=t=""}}}'''
!nawk '{t}' {d}/ain_*.conllu | LANG=C sed "{s}" >> train.txt
```

図 183: アイヌ語ミニモデル製作のための準備

^[149]安岡孝一: ローマ字・カタカナ・キリル文字併用アイヌ語 RoBERTa・DeBERTa モデルの開発, 情報処理学会研究報告, Vol.2023-CH-131 『人文科学とコンピュータ』, No.7 (2023 年 2 月 18 日), pp.1-7.

^[150]<https://ainu.ninjal.ac.jp/folklore>

^[151]<https://github.com/KoichiYasuoka/UD-Ainu>

好みに応じ、アイヌ語の他の文章を増量するのもいいでしょう。UD-Ainu の各 CoNLL-U ファイルは、train.upos にまとめ上げて、品詞付与の学習にも用います。また、これらの CoNLL-U ファイルのうち、依存構造記述が UD のフォーマットを満たしているものは、train.conllu と test.conllu に分け、係り受け解析の学習にも用います。

続けて、train.txt の各単語から、アイヌ語向けトークナイザを、my-dir/tokenizer-ain に作成しましょう。図 184 では、Unigram トークナイザ (SentencePiece^[88] の改造版) を train.txt に適用しつつ、アイヌ語のカタカナ・キリル文字をラテン文字に変換する仕掛けを組み込み、それらを DebertaV2TokenizerFast という Transformers^[26] のトークナイザに組み上げて、アイヌ語向けトークナイザを作成しています。

次に train.txt の文章から、DeBERTa モデルを my-dir/deberta-ain に作成しましょう。図 185 では、入出力幅 128 トークン・入出力ベクトル 256 次元・深さ 12 層・アテンションヘッド 4 個・中間ベクトル 768 次元という、やや小さめのモデルにしています。ただし、この DeBERTa モデルはカタカナの処理が不十分なので、単語ベクトルのコピーをおこない、新たなモデルを my-dir/deberta-ain-ext に保存しましょう (図 186)。

最後に、my-dir/deberta-ain-ext と UD-Ainu から、品詞付与・係り受け解析モデル my-dir/deberta-ain-upos を製作しましょう (図 187)。うまく my-dir/deberta-ain-upos ができたら、esupar でテストしてみましょう (図 188)。

```

from esupar.ainu import katakana,cyrillic
from transformers import DebertaV2TokenizerFast
from tokenizers import (Tokenizer,models,pre_tokenizers,normalizers,processors,
    decoders,trainers)
s=["[CLS]","[PAD]","[SEP]","[UNK]","[MASK]"]
spt=Tokenizer(models.Unigram())
spt.pre_tokenizer=pre_tokenizers.Sequence([pre_tokenizers.Whitespace(),
    pre_tokenizers.Punctuation()])
spt.normalizer=normalizers.Sequence([normalizers.Nmt()+
    [normalizers.Replace(k,v) for k,v in {**katakana,**cyrillic}.items()]+
    [normalizers.NFKD(),normalizers.Lowercase(),normalizers.StripAccents()]+
    [normalizers.Replace(i,j) for i,j in zip("bdfgjlqvwxyz","pthkcrwuhis")])
spt.post_processor=processors.TemplateProcessing(single="[CLS] $A [SEP]",
    pair="[CLS] $A [SEP] $B:1 [SEP]:1",special_tokens=[("[CLS]",0),("[SEP]",2)])
spt.decoder=decoders.WordPiece(prefix="",cleanup=True)
spt.train(trainer=trainers.UnigramTrainer(vocab_size=64000,max_piece_length=8,
    special_tokens=s,unk_token="[UNK]",n_sub_iterations=2),files=["train.txt"])
spt.save("tokenizer.json")
tkz=DebertaV2TokenizerFast(tokenizer_file="tokenizer.json",split_by_punct=True,
    do_lower_case=True,keep_accents=False,vocab_file="/dev/null")
tkz.save_pretrained("my-dir/tokenizer-ain")

```

図 184: アイヌ語向けトークナイザの作成

```

from transformers import (AutoTokenizer,DebertaV2Config,DebertaV2ForMaskedLM,
    DataCollatorForLanguageModeling,TrainingArguments,Trainer)
tkz=AutoTokenizer.from_pretrained("my-dir/tokenizer-ain",model_max_length=128)
t=tkz.convert_tokens_to_ids(["[CLS]","[PAD]","[SEP]","[UNK]","[MASK]"])
cfg=DebertaV2Config(hidden_size=256,num_hidden_layers=12,num_attention_heads=4,
    intermediate_size=768,relative_attention=True,position_biased_input=False,
    pos_att_type=["p2c","c2p"],max_position_embeddings=tkz.model_max_length,
    vocab_size=len(tkz),tokenizer_class=type(tkz).__name__,
    bos_token_id=t[0],pad_token_id=t[1],eos_token_id=t[2])
arg=TrainingArguments(report_to="none",output_dir="/tmp",save_total_limit=2,
    num_train_epochs=3,per_device_train_batch_size=64)
class ReadLineDS(object):
    def __init__(self,file,tokenizer):
        self.tokenizer=tokenizer
        with open(file,"r",encoding="utf-8") as r:
            self.lines=[s.strip() for s in r if s.strip()!=""]
        __len__=lambda self:len(self.lines)
        __getitem__=lambda self,i:self.tokenizer(self.lines[i],truncation=True,
            add_special_tokens=True,max_length=self.tokenizer.model_max_length-2)
trn=Trainer(args=arg,data_collator=DataCollatorForLanguageModeling(tkz),
    model=DebertaV2ForMaskedLM(cfg),train_dataset=ReadLineDS("train.txt",tkz))
trn.train()
trn.save_model("my-dir/deberta-ain")
tkz.save_pretrained("my-dir/deberta-ain")

```

図 185: アイヌ語ミニ DeBERTa モデルの作成

```

import torch,json
from transformers import AutoModelForMaskedLM,AutoTokenizer
mdl=AutoModelForMaskedLM.from_pretrained("my-dir/deberta-ain")
with open("my-dir/deberta-ain/tokenizer.json","r",encoding="utf-8") as r:
    spt=json.load(r)
d,c={t[0]:i for i,t in enumerate(spt["model"]["vocab"])},[]
for k,v in list(d.items()):
    s=k.replace("tu","tou").replace("ca","cia").replace("ce","cie").replace("cu",
        "ciu").replace("co","cio").replace("kk","tk").replace("pp","tp")
    if s not in d:
        c.append((k,s))
        d[s]=len(d)
e=mdl.resize_token_embeddings(len(d))
with torch.no_grad():
    for k,v in c:
        e.weight[d[v],:]=e.weight[d[k],:]
        spt["model"]["vocab"].append([v,spt["model"]["vocab"][d[k]][1]])
with open("tokenizer.json","w",encoding="utf-8") as w:
    json.dump(spt,w,ensure_ascii=False,indent=2)
tkz=AutoTokenizer.from_pretrained("my-dir/deberta-ain",use_fast=True,
    tokenizer_file="tokenizer.json")
mdl.save_pretrained("my-dir/deberta-ain-ext")
tkz.save_pretrained("my-dir/deberta-ain-ext")

```

図 186: アイヌ語向けミニ DeBERTa モデルの拡張

```

s,d="my-dir/deberta-ain-ext","my-dir/deberta-ain-upos"
!python -m esupar.train {s} {d} 32 /tmp train.upos
!python -m esupar.train {d} {d} 32 /// train.conllu test.conllu test.conllu

```

図 187: アイヌ語向け品詞付与・係り受け解析ミニモデルの製作

```

!pip install esupar
import esupar
nlp=esupar.load("my-dir/deberta-ain-upos","ainu")
doc=nlp("アイヌイタク アニ エネ アカラ ヒ エチヌカレ")
import deplacy
deplacy.serve(doc,port=None)

```

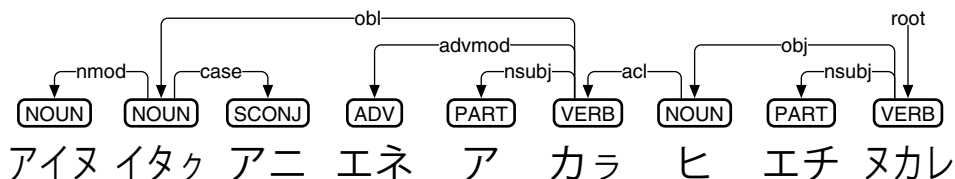


図 188: アイヌ語向け品詞付与・係り受け解析ミニモデルのテスト

Universal Dependencies と BERT/RoBERTa/DeBERTa/GPT モデルによる多言語情報処理

安岡孝一

発行日 2022 年 2 月 28 日 (初版)・2026 年 9 月 15 日 (改訂版)

発行所 京都大学人文科学研究所・未踏科学研究ユニット・

フィールドワークとデータサイエンスを融合した超学際研究ユニット

