
TeCS: A Dataset and Benchmark for Tense Consistency of Machine Translation

Yiming Ai, Zhiwei He, Kai Yu, Rui Wang

{aiyiming, zwhe.cs, kai.yu, wangrui12}@sjtu.edu.cn

Background

⌚ **FR:** Mais on les avait votés lors de la dernière période de session.

Plus-que-parfait

⌚ **EN:** But we voted on them during the last part-session.

Past simple

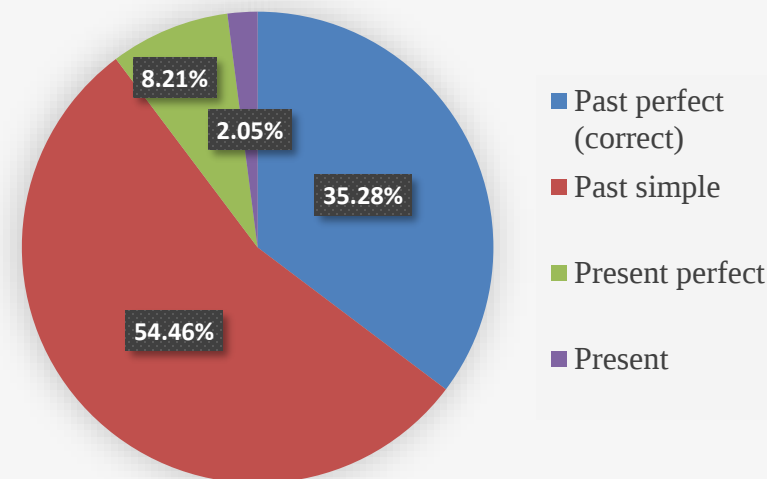
⌚ **Correction:** But we had voted on them during the last part-session.

Past perfect ✓

⌚ There are several **tense consistency errors** in the common corpora, for instance, Europarl.

⌚ **Lack of metrics** on measuring the model's mastery of tense information.

Distribution of English counterparts of French sentences in plus-que-parfait

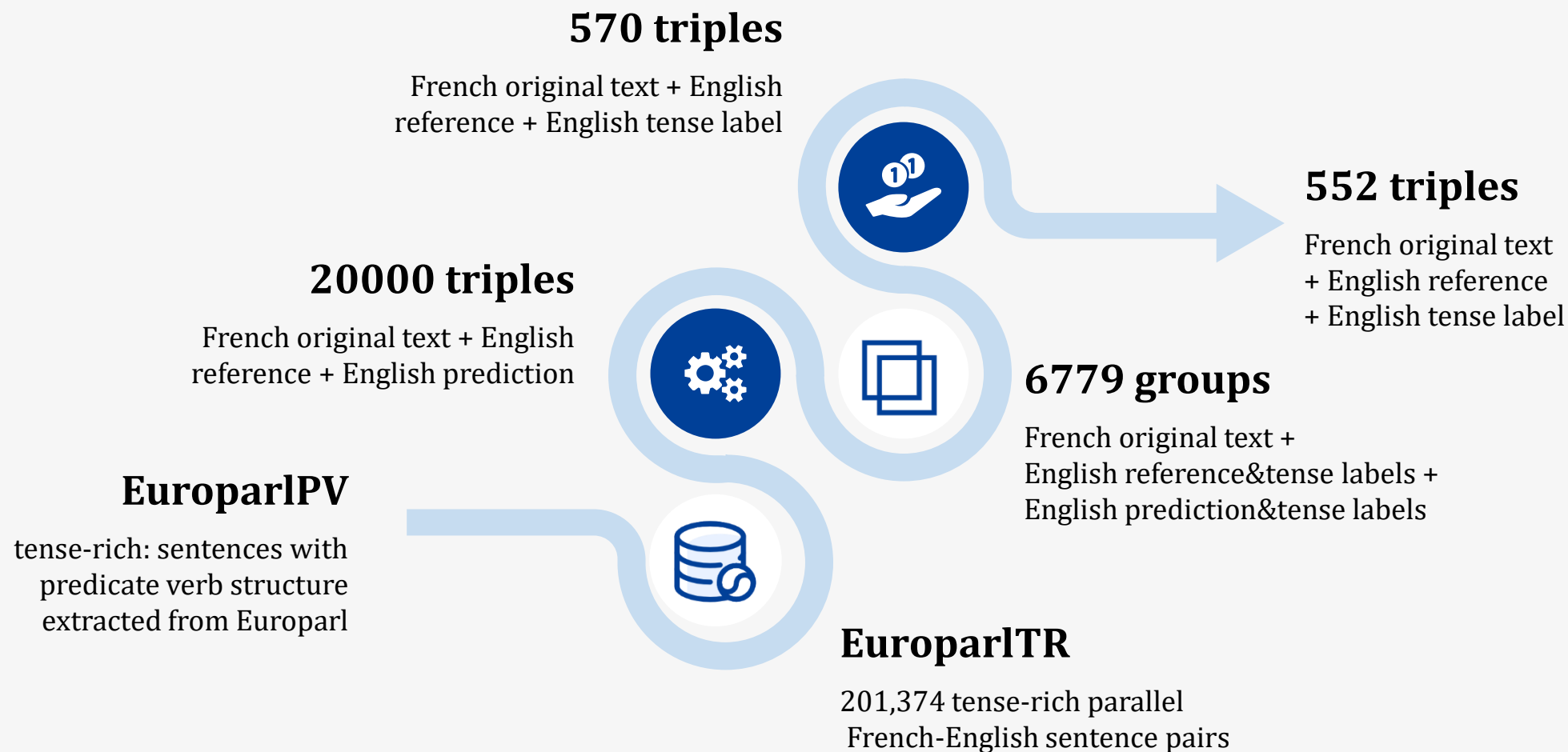


- ④ Presentation of the construction of the **tense test set**, including its tense labels
- ④ Proposal of a feasible and reproducible **benchmark** for measuring the tense consistency performance of NMT systems
- ④ **Various experiments** for different baselines with the above test set and corresponding benchmark.

Annotation Rules

- ⊗ Macro-temporal interval (**present**, **past** and **future** tenses) ✕ State of the action (**general**, **progressive** and **perfect** aspects)
- ⊗ Progressive tense $\longrightarrow$ corresponding base tense
- ⊗ Add a new category -- statements containing modal verbs

French Tenses	English Tense	Format	Example
Imparfait, Passé composé, Passé simple, Passé récent	Past simple / progressive	<i>Past</i>	That was the third point.
Présent, Future proche	Present simple / progressive	<i>Present</i>	The world is changing .
Future simple, Future proche	Future simple / progressive	<i>Future</i>	I will communicate it to the Council.
Plus-que-parfait	Past perfect	<i>PasPerfect</i>	His participation had been notified .
Passé composé	Present perfect	<i>Preperfect</i>	This phenomenon has become a major threat.
Future antérieur	Future perfect	<i>Futperfect</i>	We will have finished it at that time.
Subjonctif, Conditionnel	including Modal verbs	<i>Modal</i>	We should be less rigid.



Experiments and Conclusions

$$\text{Accuracy} = \frac{N_c}{N_t}$$

⊗ Tense (Prediction) Accuracy:

⊗ Experiments:

- **4 architectures:** Transformer, LSTM, CNN and Bi-transformer
- **3 test sets:** our tense set, Europarl test set and WMT15 test set
- **2 metrics:** BLEU and COMET
- **3 business translators:** Bing, DeepL and Google translator

⊗ Conclusions:

- By relying solely on the difference in BLEU scores on traditional test sets, we are **unable** to measure the tense prediction ability of the systems
- Our tense set can capture the **tense consistency performance**.
- To measure the tense consistency performance across different architectures, we should **focus more on tense accuracy rather than BLEU scores only**.

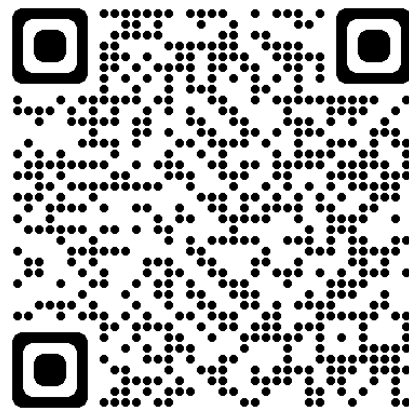
More Details?

- ④ Why adopt such annotation rules?
- ④ More details of the corpus design?
- ④ What are the characteristics of the corpus?
- ④ What specific experiments did we implement?
- ④ How about the results of the experiments?
- ④ How to use our metrics for your own experiments?

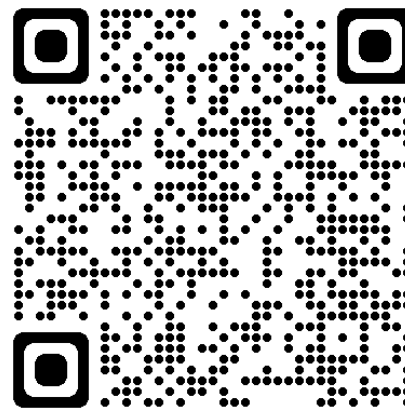
If you would like to know more:



Paper



GitHub repo



SHANGHAI JIAO TONG
UNIVERSITY





上海交通大学
SHANGHAI JIAO TONG UNIVERSITY

Thank you