



## (MACHINE) TRANSLATION EVALUATION — HUMAN EVALUATION

MULTILINGUAL NLP — LECTURE N° 4

Guillaume Wisniewski

guillaume.wisniewski@u-paris.fr

UPCité Master in Computational Linguistics

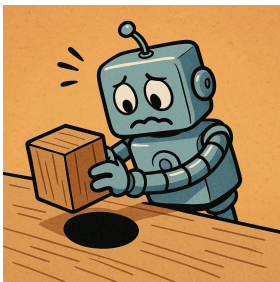
October 15<sup>th</sup>, 2025



## Part I

### WARNING

2/50



Evaluation is hard  
Evaluation is (a big part of) your job

3/50



### NO LANGUAGE LEFT BEHIND

(NLLB Team et al. 2022)



- translation models and dataset for 200 languages
- a true *tour de force*
- ↳ “Driving inclusion through the power of AI translation”
- ↳ “Translating Wikipedia for everyone”
- <https://ai.meta.com/research/no-language-left-behind/>

4/50



### SCIENCE LEFT BEHIND

“It is also a *scientifically dubious* work. In this article, I demonstrate that many of Meta AI claims made in NLLB are: unfounded, misleading, and *the result of a deeply flawed evaluation*. I will also show that, following Meta AI evaluation methodology, it is very easy to obtain even higher numbers than what they have reported.”

<http://blog.benjaminmarie.com/science-left-behind.html>

5/50



### BUT: NO SCIENCE WITHOUT MEASURING



*“When you can measure what you are speaking about, and express it in numbers, you know something about it, when you cannot express it in numbers, your knowledge is of a meager and unsatisfactory kind; it may be the beginning of knowledge, but you have scarcely, in your thoughts advanced to the stage of science.”*

Lord Kelvin, 1883

6/50



## WHY EVALUATING (MACHINE) TRANSLATION QUALITY?



### PRACTITIONERS AND USERS

- how to compare alternative systems, choices for building and customizing off-the-shelves systems?
- quantify the effectiveness of using MT

### RESEARCHERS

- evaluate models
- train our systems!

⇒ intrinsic evaluation ≠ extrinsic evaluation

7/50



## USAGE SCENARIO

### OFFLINE “BENCHMARK” TESTING OF MT ENGINE PERFORMANCE

- Typical evaluation scenario is AI/ML: reproduce a reference output
  - ↳ imitation game
- sample representative test documents with [reference human translation](#) available
- measure “similarity” between output & expected output
  - ↳ not so easy for MT

### OPERATIONAL QUALITY ASSESSMENT AT RUNTIME

- MT engine is translating new source material
- Is the output “good enough” for the underlying application?
- [reference-less evaluation](#) / confidence estimation

8/50



## COMMON USAGE OF REFERENCE-BASED EVALUATION

### ON THE SAME DATASET

- compare two MT engines
- compare two versions of the same engine
  - ↳ before and after customizing the engine
  - ↳ before and after incremental development of the engine
  - ↳ during training

### ON DIFFERENT DATASETS

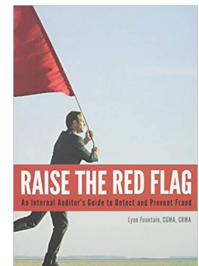
- compare MT engine performance
  - ↳ across domains or types of input data
  - ↳ on different sentence types, linguistic structures

9/50



## COMMON USAGE OF CONFIDENCE EVALUATION

- identify / flag poorly translated segment during MT engine operation
  - ↳ can the translation be used as-it?
  - ↳ is it worth correcting the translation or is it better to translate it from scratch?
- Intuitive, promising approach
  - ↳ To my knowledge: no work shows practical value
  - ↳ Useful = saves translators’ time, eases workload



10/50



## BE CAREFUL



User expectations on MT depends on their knowledge of:

- the source language
- the target language

11/50



## BUT...

Evaluation is also useful to:

- [data filtering](#): filter large-scale parallel corpora with automatic quality assessments
- [model alignment](#): optimize MT systems with quality metrics as reward/preference models
- [quality award decoding](#): develop decoding strategies that incorporate quality assessments

12/50



## AN OPEN QUESTION



- Can we use the same methods to evaluate human and machine translations?
- ↪ A natural idea
- However:
  1. Two separate communities that politely ignored each other for a long time
  2. MT evaluation developed when systems were so poor that evaluating them like humans made no sense — no longer true!
  3. Evaluation methods have been converging in recent years
- Open question: do humans and machines make the same types of errors?

13/50



## Part II

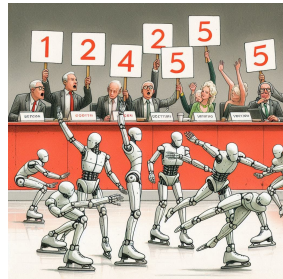
## HUMAN EVALUATION OF MT

14/50



## EARLY EVALUATION OF MT SYSTEMS

- early days of MT: manual evaluation
- Two criteria:
  - **Fluency**: how natural and grammatical the translation is in the target language.
  - **Adequacy**: how well the translation preserves the meaning of the source text.
- both evaluated on a  $[1, 5]$  scale
- used in the ALPAC Report (1966)
- formalised in the 1990s for the DARPA and NIST MT evaluations (White, O'Connell, and O'Mara 1994)



15/50



- A simple idea, a complex implementation
- Bilingual or non-bilingual annotators
- Providing a reference translation
- Defining appropriate scales
- Better to rank translation hypothesis or to score them?
- ...

16/50



## FLUENCY AND ADEQUACY SCALES

| Adequacy |                |
|----------|----------------|
| 1        | no meaning     |
| 2        | little meaning |
| 3        | much meaning   |
| 4        | most meaning   |
| 5        | all meaning    |

| Fluency |                    |
|---------|--------------------|
| 1       | incomprehensible   |
| 2       | disfluent English  |
| 3       | non-native English |
| 4       | good English       |
| 5       | flawless English   |

17/50



## AN IMPORTANT — AND OFTEN OVERLOOKED — QUESTION



### WHO PRODUCES THE ANNOTATIONS?

- MT experts
- (Professional) translators
- Crowd-workers

⇒ huge impact on evaluation quality (more on that latter)

18/50



## WITH A PICTURE:

| Judge Sentence                                                                                         |                        |                        |
|--------------------------------------------------------------------------------------------------------|------------------------|------------------------|
| You have already judged 14 of 3064 sentences, taking 86.4 seconds per sentence.                        |                        |                        |
| Source: les deux pays constituent plutôt un laboratoire nécessaire au fonctionnement interne de l'ue . |                        |                        |
| Reference: rather, the two countries form a laboratory needed for the internal working of the eu .     |                        |                        |
| Translation                                                                                            | Adaptivity             | Fluency                |
| both countries are rather a necessary laboratory the internal operation of the eu .                    | C C C C R<br>1 2 3 4 5 | C C C C R<br>1 2 3 4 5 |
| both countries are a necessary laboratory at internal functioning of the eu .                          | C C R C C<br>1 2 3 4 5 | C C R C C<br>1 2 3 4 5 |
| the two countries are rather a laboratory necessary for the internal workings of the eu .              | C C C R C<br>1 2 3 4 5 | C C C R C<br>1 2 3 4 5 |
| the two countries are rather a laboratory for the internal workings of the eu .                        | C C R C C<br>1 2 3 4 5 | C C R C C<br>1 2 3 4 5 |
| the two countries are rather a necessary laboratory internal workings of the eu .                      | C C R C C<br>1 2 3 4 5 | C C R C C<br>1 2 3 4 5 |
| Annotator: Philipp Koehn Tasks: WMT06 French-English                                                   |                        | Annotate               |
| Instructions                                                                                           | Se All Meaning         | Se Fluent English      |
|                                                                                                        | Se Most Meaning        | Se Good English        |
|                                                                                                        | Se Much Meaning        | Se Non-native English  |
|                                                                                                        | Se Little Meaning      | Se Different English   |
|                                                                                                        | Se None                | Se Incomprehensible    |

19/50



## LET'S TRY IT!

- Source: N'y aurait-il pas comme une vague hypocrisie de votre part?
- Reference: Is there not an element of hypocrisy on your part?
- System 1: Would it not as a wave of hypocrisy on your part?
- System 2: Is there would be no hypocrisy like a wave of your hand?
- System 3: Is there not as a wave of hypocrisy from you?

20/50



## MEASURING ANNOTATOR AGREEMENT

### THE PROBLEM

- Agreement: the proportion of annotations on which two annotators agree.
- ↳ "Agree" means they give the same label to the same item.
- Problem (from a statistician's point of view): even if two people label completely at random, they will sometimes agree by chance.
- ↳ We therefore need a more precise measure of agreement.

21/50



## MEASURING ANNOTATOR AGREEMENT

### COHEN'S $\kappa$

- Quantifies the extent to which two annotators agree on how a set of data items should be labelled.
  - ↳ Pre-defined set of labels  $\Phi$  set of items to be annotated.
  - ↳ Each item is labelled independently by both annotators.
- Measures the proportion of items assigned the same label by both annotators.
- Adjusts for the fact that some level of agreement may occur purely by chance.

$$\kappa = \frac{P_a - P_e}{1 - P_e}$$

- ↳  $P_a$ : observed proportion of agreement (actual agreement).
- ↳  $P_e$ : expected proportion of agreement by chance.

21/50



## THE LANDIS AND KOCH 1977 INTERPRETATION SCALE

| $\kappa$  | interpretation |
|-----------|----------------|
| < 0.00    | Poor agreement |
| 0.00–0.20 | Slight         |
| 0.21–0.40 | Fair           |
| 0.41–0.60 | Moderate       |
| 0.61–0.80 | Substantial    |
| 0.81–1.00 | Almost perfect |

- The scale everyone uses — basically because it's the only one that exists.
- In short: "good enough" when  $\kappa > 0.5$ ... if we're feeling optimistic.

22/50



## EXAMPLE (1)

- First annotator:
  - ["B", "B", "B", "A", "C", "C", "B", "A", "A", "C", "C", "B", "A", "C", "C", "C", "A"]
- Second annotator:
  - ["B", "B", "C", "A", "C", "C", "C", "A", "A", "B", "C", "B", "A", "C", "C", "C", "A"]

23/50



## EXAMPLE (2)

Contingency table:

|   | A | B | C |
|---|---|---|---|
| A | 5 | 1 | 0 |
| B | 0 | 4 | 2 |
| C | 0 | 2 | 6 |

↪ how often the two annotators assigned the same or different labels to the same items.

24/50



## EXAMPLE (3)

$$P_a = \frac{5 + 4 + 6}{5 + 1 + 4 + 2 + 2 + 5} = \frac{15}{20} = 0.75$$

$$P_e = \sum_{c \in \{A, B, C\}} p_1(c) \times p_2(c)$$

After a rather tedious round of calculations:

$$P_e = 0.34$$

$$\kappa \approx 0.62$$

25/50



## EXTENSIONS & CHALLENGES

### VARIANTS OF AGREEMENT MEASURES

- **Weighted  $\kappa$** : takes into account the degree of disagreement.
- **Fleiss'  $\kappa$** : generalises Cohen's  $\kappa$  to multiple annotators.
- **ICC (Intraclass Correlation Coefficient)**: suitable for continuous or ordinal data.

### LIMITATIONS AND PRACTICAL ISSUES

- **High cost**: requires multiple annotations per item, which may reduce the total number of annotated items.
- **Interpretation can be difficult**: what level of  $\kappa$  is "good enough" depends on the context.
- **What to do when results are poor?** (e.g., low agreement)
  - Improve annotation guidelines or training.
  - Refine label definitions.
  - Reconsider task complexity.

26/50



## CONSISTENCY OF JUDGMENTS OF MT

all figures are coming from WMT Evaluation Campaign reports

| Inter-rater agreement |          |
|-----------------------|----------|
|                       | $\kappa$ |
| 2011                  | 0.40     |
| 2012                  | 0.33     |
| 2013                  | 0.26     |
| 2014                  | 0.37     |

| Intra-rater agreement |          |
|-----------------------|----------|
|                       | $\kappa$ |
| 2011                  | 0.58     |
| 2012                  | 0.41     |
| 2013                  | 0.48     |
| 2014                  | 0.52     |

↪ very low in both cases

↪ Would the result be the same if the translations were better?

27/50



## MORE BROADLY: LIMITATIONS OF ADEQUACY & FLUENCY EVALUATION

1. Subjectivity and inter-annotator variance
  - ↪ annotators interpret criteria differently
2. Scale bias
  - ↪ some annotators are stricter or more generous
3. High correlation between fluency and adequacy
  - ↪ redundant, inefficient measures
4. Poor sensitivity to small differences
  - ↪ discrete scales fail to capture fine quality gaps
  - ↪ unstable system rankings

28/50



## WHY "GOOD TRANSLATION" IS A PROBLEMATIC NOTION

- There is no single, universal definition of translation quality.
- **Quality can refer to:**
  - **Semantic**: preservation of meaning.
  - **Stylistic**: tone and register appropriateness.
  - **Functional**: effectiveness in achieving communicative goals.
- What counts as "good" depends on:
  - The **purpose** of the translation (legal, literary, technical...).
  - The **target audience** and its expectations.
  - The **socio-cultural context**.

⇒ Quality in translation is **context-dependent**, not absolute.



29/50



## TRANSLATION STUDIES: FUNCTION OVER FORM

- From linguistic equivalence to **functional approaches**.
- **Skopos theory** (Vermeer 1978):
  - A translation is “good” if it fulfils its communicative function.
  - Multiple translations can be equally valid.
- Other influential frameworks:
  - Reiss and Vermeer 1984 — text type and function.
  - House 2015 — functional-pragmatic equivalence.

### KEY IDEA

Quality ≠ fidelity to one “true” source text; it is evaluated against **purpose, context and audience expectations**.

30/50



## Part III

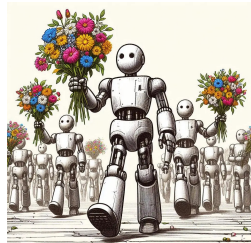
## HUMAN EVALUATION IN THE 21ST CENTURY

31/50



## DIRECT ASSESSMENT (DA)

- introduced by Graham, Baldwin, Moffat, et al. 2013 and Graham, Baldwin, and Mathur 2015
- one simple, “global” question  
*To what extent does this translation adequately convey the meaning of the source text?*
- ↪ combines fluency and adequacy
- continuous evaluation on a 1–100 scale
- ↪ the annotator moves a slider to a value reflecting the perceived quality
- collect several independent judgements for each segment (typically 3 to 5)
- ↪ **aggregate scores (median or mean)**



The Revolution of Flowers (in MT)

32/50



## TWO CORE IDEAS

### KISS



- keep it simple — no overcomplicating things
- give annotators as much freedom as possible

### THE UNREASONABLE EFFECTIVENESS OF DATA

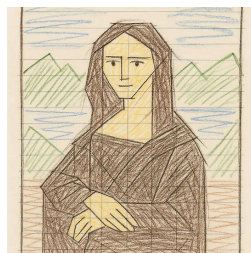
- Alon Halevy, Peter Norvig, and Fernando Pereira
- more data > trying to control data quality
- simple idea: law of large numbers
- ↪ the average of many ratings = a better estimate of the true score

33/50



## BECAUSE THINGS SHOULDN'T BE TOO SIMPLE

- Control annotation quality with:
  1. “trap” sentences (where the correct translation is obvious)
  2. perfect reference sentences inserted into the sample
- normalise scores per annotator to correct individual scale biases
- ↪ typically, a z-score is sufficient



34/50



## WHY DA IS THE ANSWER TO EVERYTHING

if only E. M. had studied machine translation...

1. Scale bias correction
  - ↪ scores are normalised per annotator (z-scoring)
  - ↪ neutralises strict vs. generous raters
  - ↪ enables cross-annotator comparability
  - ↪ improves aggregate robustness
2. Less inter-annotator variance
  - ↪ combining fluency and adequacy into one intuitive rating reduces criterion divergence
  - ↪ significant increase in  $\kappa$  ( $\geq 0.6$ )
3. Continuous scale  $\Rightarrow$  higher sensitivity
  - ↪ captures subtle quality differences
  - ↪ allows statistical smoothing
4. Alignment with real human perception
  - ↪ DA simplifies evaluation  $\rightarrow$  one global, intuitive score

35/50

- no information about the nature of errors
- très souvent complété par
  - des annotations d'erreurs (MQM),
  - des comparaisons par paires



36/50

## Part IV

### DIAGNOSTIC-BASED EVALUATION

37/50

- How are human translations evaluated?
- ↳ But surely: humans don't make mistakes!
- Core idea: annotate errors
- ↳ identify the error span (words involved)
- ↳ label describing the error type
- No numerical evaluation
- ↳ not of interest to translators!
- ↳ goal = understand errors to improve

Établissements de traduction de la parole vers le texte \*\*LA-TL-INS, LA-TL-INS\*\* pour les humains humains \*\*LA-TL-INS, LA-TL-INS\*\*  
 Nous présentons SUT, le système intégré de traduction de la parole vers le texte \*\*LA-TL-INS, LA-TL-INS\*\* basé sur POCKETSPHINX et MOSES. Il est composé à différentes lignes de base \*\*LA-TL-INS, TR-SG-TL\*\* basées sur \*\*LA-TL-INS, LA-TL-INS\*\* ANIS — le système de transcription de la parole vers le texte \*\*LA-TL-INS, LA-TL-INS\*\* développé par le groupe Speech de LORIA, MOSES et les outils de traduction de Google. Un petit corpus de transcriptions de référence de français isolées \*\*LA-TL-INS, LA-TL-INS\*\* provenant de la campagne d'évaluation (ESTER) a été traduit par des experts humains pour l'évaluation. Le jeu d'essai des mots (W2T) est utilisé pour reconnaître des données synthétiques en français et BEU est utilisé pour tester \*\*LA-TL-INS, TR-TL-TE, LA-TL-INS\*\* les traductions. En outre, les transcriptions de référence sont traduites automatiquement à l'aide de \*\*LA-TL-INS\*\* MOSES et de GOOGLE afin d'évaluer l'impact des erreurs de reconnaissance sur la qualité de la traduction.

38/50

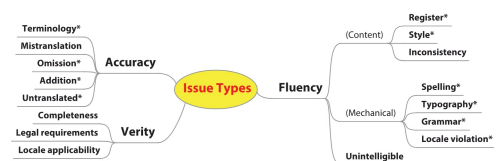
- List of possible errors
- Often organised hierarchically

39/50

#### PRINCIPLE

- Annotators do not give a global score,
  - ↳ they identify and classify errors in the translation.
  - Each error is annotated with:
    1. a category (e.g. Accuracy, Fluency, Terminology...)
    2. a more specific subtype (e.g. Omission, Mistranslation, Grammar...)
    3. a severity level (Minor / Major / Critical)
    4. the exact span in the translation (and often in the source)
  - Annotations are then aggregated to calculate scores, often weighted by severity.
  - (Lommel, Burchardt, and Uszkoreit 2014)
- ⇒ A diagnostic rather than purely evaluative approach.

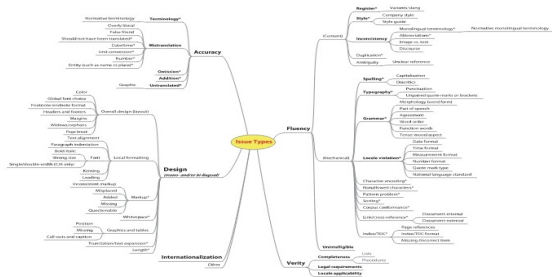
40/50



41/50



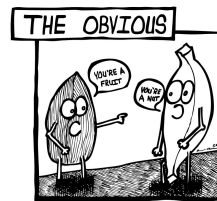
## THE MQM (FULL) TYPOLOGY



42/50



## STATING THE OBVIOUS



- Error annotation requires experts trained in (i) translation and (ii) the use of this typology
- Time-consuming

43/50



## COST AND COMPARISON WITH DIRECT ASSESSMENT

|                    | DA                            | MQM                                         |
|--------------------|-------------------------------|---------------------------------------------|
| Type of judgement  | Global score (0–100)          | Detailed error annotation                   |
| Time per segment   | 15–30 s                       | 2–5 min (or more for longer sentences)      |
| Annotator training | Simple                        | Complex (categories, severity, calibration) |
| Analysis           | Statistical (means, z-scores) | Qualitative + quantitative                  |
| Granularity        | Global                        | Fine (specific error categories)            |
| Cost               | Low to medium                 | High                                        |

44/50



## EXAMPLE OF ANNOTATION WITH MQM

• Filters (17221 of 17221 rows pass filters) [Clear all filters](#) ☐ Select only the segments for which all or no systems have scores

| Sample segments          | Doc | Design | System          | Source                                            | Target                                                                                                                                                                                                                      | Severity       | Category                                         | Score |
|--------------------------|-----|--------|-----------------|---------------------------------------------------|-----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|----------------|--------------------------------------------------|-------|
| citizenews.com 22 1 1454 |     |        | hpa-ModelCline  | 在当地时间下午，一架小型飞机在德意志联邦共和国首都柏林上空坠毁，造成多人伤亡。事故原因正在调查中。 | On the afternoon of the 20th local time, a small plane crashed into the top floor of a multi-story residential building in West <b>Germany</b> , North Rhine-Westphalia, Germany, killing three people in the accident.     | <b>CRASHED</b> | Major locale convention/address for road context | 100   |
| citizenews.com 22 1 1454 |     |        | hpa-0001-NEP    | 在当地时间下午，一架小型飞机在德意志联邦共和国首都柏林上空坠毁，造成多人伤亡。事故原因正在调查中。 | On the afternoon of the 20th local time, a small plane crashed into the top floor of a multi-story residential building in West <b>Germany</b> , North Rhine-Westphalia, Germany. Three people were killed in the accident. | <b>CRASHED</b> | Major locution/AMETion context                   | 100   |
| citizenews.com 22 1 1454 |     |        | hpa-Facebook-41 | 在当地时间下午，一架小型飞机在德意志联邦共和国首都柏林上空坠毁，造成多人伤亡。事故原因正在调查中。 | On the afternoon of the 20th local time, a small plane crashed into the top floor of a multi-story residential building in West <b>Germany</b> , North Rhine-Westphalia, Germany, killing a total of three people.          | <b>CRASHED</b> | Major locution/AMETion context                   | 100   |
| citizenews.com 22 1 1454 |     |        | hpa-110-ME      | 在当地时间下午，一架小型飞机在德意志联邦共和国首都柏林上空坠毁，造成多人伤亡。事故原因正在调查中。 | On the afternoon of the 20th local time, a small plane crashed into the top floor of a multi-story residential building in West <b>Germany</b> , North Rhine-Westphalia, Germany. Three people were killed in the accident. | <b>CRASHED</b> | Major locution/AMETion context                   | 100   |

45/50



## A LARGE-SCALE STUDY OF HUMAN EVALUATION FOR MACHINE TRANSLATION

(Freitag et al. 2021)

### RESEARCH QUESTION

- Observation: increasingly difficult to evaluate MT output
  - higher quality ⇒ errors are finer and subtler
- Core idea: all evaluation relies on error identification
  - error-based evaluation
  - evaluate evaluation using MQM

### METHODOLOGY

- Manual MQM annotation of WMT'20 English→German and Chinese→English system outputs
  - largest MQM dataset ever collected
- Compare expert MQM rankings with crowd rankings (WMT) and automatic metrics

46/50



## MAIN FINDINGS

1. Expert MQM rankings can substantially differ from traditional crowd rankings (notably restoring a preference for human translation).
2. Some modern automatic metrics (especially embedding-based) align more closely with expert judgements than crowd evaluations.
3. Public release of a large MQM-annotated corpus to support comparison, supervised metric training, and further research.
4. In-depth analysis of error type distribution, inter-annotator disparities, contextual effects, and more.

47/50



## Part V

### CONCLUSION

48/50



## THE 3 STEPS OF HUMAN EVALUATION IN MT

### OVERVIEW OF EVOLUTION

Three steps:

1. 1990 → 2015: fluency & adequacy  
↳ collected in various forms
2. 2015 → ... DA
3. 2020 → ... MQM

### AN IMPORTANT NOTE

- same evolution in different (sub)fields of NLP
- increasing expertise in translation

49/50



## Part VI

### REFERENCES

50/50



- Freitag, Markus et al. (2021). "Experts, Errors, and Context: A Large-Scale Study of Human Evaluation for Machine Translation". In: *Transactions of the Association for Computational Linguistics* 9. Ed. by Brian Roark and Ani Nenkova, pp. 1460–1474.
- Graham, Yvette, Timothy Baldwin, and Nitika Mathur (May 2015). "Accurate Evaluation of Segment-level Machine Translation Metrics". In: *Proceedings of the 2015 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies*. Ed. by Rada Mihalcea, Joyce Chai, and Anoop Sarkar. Denver, Colorado: Association for Computational Linguistics, pp. 1183–1191.
- Graham, Yvette, Timothy Baldwin, Alistair Moffat, et al. (Aug. 2013). "Continuous Measurement Scales in Human Evaluation of Machine Translation". In: *Proceedings of the 7th Linguistic Annotation Workshop and Interoperability with Discourse*. Ed. by Antonio Pareja-Lora, Maria Liakata, and Stefanie Dipper. Sofia, Bulgaria: Association for Computational Linguistics, pp. 33–41.
- House, Juliane (2015). *Translation Quality Assessment: Past and Present*. Routledge.

50/50



- Landis, J. Richard and Gary G. Koch (1977). "The Measurement of Observer Agreement for Categorical Data". In: *Biometrics* 33.1, pp. 159–174.
- Lommel, Arle Richard, Aljoscha Burchardt, and Hans Uszkoreit (Dec. 2014). "Multidimensional Quality Metrics (MQM): A Framework for Declaring and Describing Translation Quality Metrics". In: *Tradumàtica: tecnologies de la traducció* 12. Ed. by Attila Görög and Pilar Sánchez-Gijón, pp. 455–463.
- NLLB Team et al. (2022). *No Language Left Behind: Scaling Human-Centered Machine Translation*.
- Reiss, Katharina and Hans J. Vermeer (1984). *Grundlegung einer allgemeinen Translationstheorie*. Niemeyer.
- Vermeer, Hans J. (1978). "Ein Rahmen für eine allgemeine Translationstheorie". In: *Lebende Sprachen* 23.3, pp. 99–102.
- White, John S., Theresa A. O'Connell, and Francis E. O'Mara (Oct. 1994). "The ARPA MT Evaluation Methodologies: Evolution, Lessons, and Future Approaches". In:

50/50



*Proceedings of the First Conference of the Association for Machine Translation in the Americas*. Columbia, Maryland, USA.

50/50