

Comparing UMR and Cross-lingual Adaptations of AMR

Shira Wein

Georgetown University
sw1158@georgetown.edu

Julia Bonn

University of Colorado at Boulder
julia.bonn@colorado.edu

Abstract

Abstract Meaning Representation (AMR) is a popular semantic annotation schema that presents sentence meaning as a graph while abstracting away from syntax. It was originally designed for English, but has since been extended to a variety of non-English versions. These cross-lingual adaptations, to varying degrees, incorporate language-specific features necessary to effectively capture the semantics of the language being annotated. Uniform Meaning Representation (UMR) on the other hand, the multilingual extension of AMR, was designed specifically for uniform cross-lingual application. In this work, we discuss these two approaches to extending AMR beyond English. We describe both approaches, compare the information they capture for a case language (Spanish), and outline implications for future work.

1 Introduction

Abstract Meaning Representation (AMR; [Banarescu et al., 2013](#)) is a symbolic meaning representation which captures the meaning of a sentence in the form of a directed, rooted graph composed of predicate argument structures. AMR was originally designed for English, but has since been extended to many other languages. These cross-lingual adaptations of AMR vary in their approach to adapting English-centric AMR to other languages, which has posed a number of challenges.

In addition to language- or language family-specific ([Heinecke and Shimorina, 2022](#)) adaptations of AMR, Uniform Meaning Representation (UMR; [Van Gysel et al., 2021a](#)) is a recent multilingual extension of AMR which attempts to be generally cross-lingually portable.

Approach to cross-lingual adaptation has a significant impact on the utility of the annotated data. Formalisms which have similarly structured parallel annotations are better suited for incorporation

Sentence: *He denied any wrongdoing.*

AMR:

```
(d / deny-01
  :ARG0 (h / he)
  :ARG1 (w / wrong-02
    :mod (a / any)
    :ARG0 h))
```

UMR:

```
(s1d / deny-01
  :ARG0 (s1p / person
    :ref-person 3rd
    :ref-number Singular)
  :ARG1 (s1t/ thing
    :ARG1-of (s1d2/ do-02
      :ARG0 s1p
      :ARG1-of (s1w/ wrong-02)
      :MODPRED s1d))
  :ASPECT Performance
  :MODSTR FullAff)

(s1 / sentence
  :temporal ((DCT :before s1d)
    (s1d :before s1d2))
  :modal ((AUTH :FullAff s1p)
    (s1p :FullAff s1d)
    (s1d :Unsp s1d2))
  :coref (s0p :same-entity s1p))
```

Figure 1: AMR and UMR (from the guidelines, <https://github.com/umr4nlp/umr-guidelines/blob/master/guidelines.md>) annotating the same sentence.

into downstream applications, such as structure-aware machine translation systems ([Sulem et al., 2015](#)). Therefore, it is critical to understand the differences between UMR and cross-lingual adaptations of AMR, with regard to what linguistic information they encode, as it will impact the functionality of the annotations.

Though strong efforts have been made to adapt AMR to cross-lingual contexts in two directions (individual cross-lingual AMR extensions, and the more expansive UMR), there has not yet been any comparison between the effectiveness and comprehensiveness of these two different approaches.

In this work, we examine differences between these attempts at fashioning non-English-centric versions of AMR. In §2, we outline cross-lingual adaptations of AMR, including both annotation schema and generation/parsing tools, and survey select adaptations. Next (§3), we introduce UMR, the multilingual extension of AMR. In §4, we take a close look at how UMR and a cross-lingual adaptation of AMR handle linguistic features and examine cross-lingual challenges to UMR/AMR annotation. Finally, in §5, we discuss challenges for both UMR and cross-lingual extensions of AMR.

2 Cross-lingual Adaptations of AMR

AMR is designed to abstract away from the surface-form and syntactic nuance of the sentence, focusing only the basic meaning. In AMR annotations, nodes reflect concepts and the edges are labeled with relations between the concepts. Annotation of AMR concepts relies in part on PropBank lexicon of frame files¹ (Kingsbury and Palmer, 2002; Palmer et al., 2005; Pradhan et al., 2022), by annotating the frame associated with the token as the concept in the AMR graph.

Though AMR was designed exclusively for English and was not intended to be an interlingua (Banarescu et al., 2013), it has now been extended to multiple languages. Table 1 contains the cross-lingual AMR adaptations to date, with their publications as well as the underlying resources (frame files) they use and the corpus they annotate.

AMR has also been assessed as an interlingua for Czech (Urešová et al., 2014), Chinese (Xue et al., 2014; Wein et al., 2022b), and Spanish (Wein and Schneider, 2021). Xue et al. (2014) explores the adaptability of English AMR to Czech and Chinese. The authors suggest that, although it was not designed to be an interlingua, AMR may be cross-linguistically adaptable because it abstracts away from morphosyntactic differences. Cross-linguistic comparisons between English/Czech and English/Chinese AMR pairs indicate that most pairs align well, though there are some instances of divergence due to insertions, for example.

¹<https://github.com/propbank>

Urešová et al. (2014) describes the types of differences between AMRs for parallel English and Czech sentences, and finds that the differences may be either due to convention/surface-level nuances which could be changed in the annotation guidelines, or may be due to inherent facets of the AMR annotation schema. One notable area of difference stems from the appearance of language-specific idioms and phrases.

Recent work has defined the types and causes of divergences between cross-lingual AMR pairs for English-Spanish parallel sentences. The causes of structural differences between parallel AMRs are identified as being due to semantic divergences, syntactic divergences, or annotation choices (Wein and Schneider, 2021).

In the subsections that follow, we consider four adaptations of AMR to individual languages.

2.1 Chinese AMR Adaptation

Li et al. (2016) suggested that AMR would be particularly well adapted to languages which vary morphosyntactically from English, because AMR abstracts away from the surface syntactic structure, motivating adaptation to Chinese. The Chinese AMR (CAMR) annotation schema largely matches that of the English annotation schema, with the concepts being tokens in Chinese instead of English. Notably, Chinese has very little inflectional morphology, so the AMR concepts more often directly correspond to tokens in the sentence than in English annotation. Extensions to the annotation guidelines are made for Chinese-specific constructions, including but not limited to (1) number and classifier construction, (2) serial-verb construction, (3) headless relative construction, (4) verb-complement construction, (5) split verb construction, and (6) reduplication. In the case where reduplication signals intensified meaning, Chinese AMR annotates this with another abstract concept, often with the role :UNIT. Discourse relations are also represented with concepts from the Chinese Discourse Treebank (DCTB; Zhou and Xue, 2015). These adaptations to the guidelines were identified during the annotation process.

2.2 Portuguese AMR Adaptations

Two distinct Portuguese AMR annotation schemata have been developed. Anchiêta and Pardo (2018) annotated the Portuguese translation of *The Little Prince*, and aligned the Portuguese sentences with the English ones (though there is one more sen-

Language	Underlying Resource(s)	Corpus	Publication
English	English PropBank	The Little Prince	Banarescu et al. (2013)
Chinese	Chinese Discourse Treebank	The Little Prince	Li et al. (2016)
Spanish	English PropBank	The Little Prince	Miguelles-Abraira et al. (2018)
Spanish	AnCora	AMR 3.0 Data (news etc.)	Wein et al. (2022a)
Portuguese	FrameSet Verbo-Brasil	The Little Prince	Anchieta and Pardo (2018)
Portuguese	FrameSet Verbo-Brasil	News, PropBank.Br	Sobrevilla Cabezedo and Pardo (2019)
Vietnamese	Vietnamese comp. lexicon	The Little Prince	Linh and Nguyen (2019)
Korean	Korean PropBank	ExoBrain	Choe et al. (2020)
Turkish	[Unspecified]	The Little Prince	Azin and Eryigit (2019)
Turkish	Turkish PropBank	The Little Prince	Oral et al. (2022)
Persian	Perspred, English PropBank	The Little Prince	Takhshid et al. (2022)

Table 1: Comparison of characteristics of the AMR cross-lingual adaptations. “Underlying Resource(s)” for AMR reflect the lexicon or frameset used to mark roles and senses of concepts. “Corpus” indicates the corpus selected for annotation of the schema.

tence in the Portuguese corpus). This approach to Portuguese AMR annotation consists of importing the English AMR annotation for the aligned sentences, and changing the PropBank concepts to the equivalent Portuguese concepts from Frameset Verbo-Brasil (Sanches Duran and Aluísio, 2015). Any linguistic features that cause Portuguese AMR annotation to differ structurally from English AMR annotation were adjudicated upon at time of annotation for a given sentence. For example, instances of implied subjects and the particle “se”.

A second Portuguese AMR annotation schema was developed shortly afterwards, which translates and fully adapts the English AMR guidelines to Portuguese. Duran and Aluísio (2011) annotated news texts from the *Folha de São Paulo* Brazilian news agency and from the *PropBank.Br* corpus. The verb senses are again determined by framesets from Verbo-Brasil. Modal verbs, which do not appear in Verbo-Brasil, are replaced by their direct Portuguese translations. Linguistic features handled specially in these new Portuguese AMR guidelines include use of the 3rd person singular and indeterminate subjects. Notably, multi-word expressions are replaced by their nearest one-word synonym.

2.3 Vietnamese AMR Adaptation

When adapting AMR to Vietnamese (Linh and Nguyen, 2019), the focus was on demonstrating relationships between entities and expanding annotation to include labels that mark function words, tense, and gender. Concepts were mapped from English to Vietnamese using the Vietnamese computational lexicon (Nguyen et al., 2006), with the addition of some new concepts. Linguistic differences between English and Vietnamese that trigger different annotation include morphosyntactic real-

ization of manner as well as the presence of noun classifiers in Vietnamese. In English, manner is frequently expressed through *-ly* adverbs. In English AMR, *-ly* adverbs aren’t included in graphs; rather, they are replaced by a related roleset or a related nominal or adjectival concept under a :MANNER relation (e.g. *quickly* in the surface form becomes :MANNER (q / quick) in the graph), Vietnamese expresses manner adjectivally, so such adjustments are unnecessary. In Vietnamese AMR, noun classifiers are omitted from the representation, except in cases where a noun classifier is alone (not directly preceding a noun phrase). Here, the co-referent needs to be included in the graph.

2.4 Korean AMR Adaptations

Choe et al. (2019) establishes a desire to make a Korean AMR annotation as similar as possible to AMR annotation in other languages so that cross-lingual annotations will be compatible and comparable, while at the same time bolstering the schema’s ability to accurately reflect Korean semantics. The main areas in which special adaptations were needed include the copula and its negation, as well as case-stacking where multiple subjects or objects are involved.

Choe et al. (2020) further develops the annotation schema for Korean AMR and releases an annotated corpus for texts using Korean PropBank frames. Annotations were piloted on the ExoBrain Corpus, the Korean translation of The Little Prince, and example sentences for verbs in the Basic Korean Dictionary; the actual released corpus consists of annotations on the ExoBrain Corpus. The abstract rolesets used in English AMR (such as have-org-role-91) are also used for Korean AMR. For copular annotation, the use of :domain and :polarity are expanded.

3 UMR

The recent development of the Uniform Meaning Representation (Van Gysel et al., 2021a) aims to incorporate uniform treatments for linguistic diversity into the AMR annotation process.

Uniform Meaning Representation (UMR) is designed to extend AMR to a cross-linguistically viable meaning representation. Related work on BabelNet Meaning Representation (Navigli et al., 2022; Martínez Lorenzo et al., 2022) also extends AMR to a multilingual context, by moving away from English PropBank and instead using VerbAtlas (Di Fabio et al., 2019) for cross-lingual frames and BabelNet concept inventory (Navigli and Ponzetto, 2010).

To accommodate cross-linguistic diversity, UMR incorporates paradigmatic lattices to organize annotation categories from coarse-grained to more specific. Annotators are able to use the degree of granularity that is most suitable for the grammar of the language being annotated. Lattices produced for this purpose indicate degrees of granularity for discourse relations, modality, number, spatial relations, aspect, and temporality. The number of concepts associated with any given token (polysynthesis and agglutination) can also vary by language, so UMR does not require that morphologically complex words be broken down into separate morphemes when being annotated as concepts—however, it builds in the ability to do so where appropriate to support uniformity.

UMR extends AMR in 3 core ways: (1) it is capable of annotating low-resource languages, (2) it more comprehensively annotates modality, aspect, quantification, and scope for the benefit of logical inference, and (3) it annotates temporal, modal, and coreference relations across sentences.

At the sentence level, UMR adds aspect, modal strength, and quantifier scope attribute roles. Aspect is annotated for events and states at five base level values, with finer-grained values in lattice format (e.g., :ASPECT STATE). Sentence-level modal annotation comes in three strengths for both affirmative and negative (e.g., :MODSTR PRTAFF for partial-affirmative). The optional scope node augments predicates.

At the document level, UMR adds temporal and modal dependencies, plus coreference. Document-level semantic relations can be created for concepts/events within a sentence or across sentence boundaries. These document-level relations are

able to be more fine-grained and provide more detailed information than their sentence-level counterparts, for instance, document-level modal relations are able to mark a conceiver in addition to the strength and polarity marked at the sentence level.

While UMR follows AMR in using existing role-set lexicons where possible (referred to as Stage 1 annotation), languages without these resources can also be annotated in UMR (Stage 0 annotation). During Stage 0 annotation, UMR-Writer (Zhao et al., 2021) allows annotators to select tokens for use as graph predicates and then add those predicates into a lexicon. Argument structures for these predicates are added using UMR’s inventory of participant and non-participant roles. The predicates added to the working lexicon in combination with their participant role annotation information can be used to generate a roleset lexicon, moving a language from Stage 0 to Stage 1 annotation.

Recent work on UMR has produced small sets of annotations for four indigenous languages (Kukama, Arapaho, Sanapaná, and Navajo) (Van Gysel et al., 2021b), an online application (UMR-Writer) for producing AMR annotations (Zhao et al., 2021), automatically annotating tense and aspect in UMR (Chen et al., 2021), and incorporating non-verbal interactions into UMR annotation (Lai et al., 2021). Bonn et al. (2023) outlined deterministic conversion of AMRs to UMRs, specifically the roles, rolesets, and concepts.

4 Differences Between UMR and Cross-lingual AMR

In this section, we compare the specific linguistic features that both schemata encode, and consider two noteworthy obstacles/factors to successful annotation of UMR *or* AMR: idiomatic phrases and reliance on English concepts.

4.1 Comparison with Spanish AMR

In order to perform a language-specific comparison between a cross-lingual extension of AMR and UMR, we compare what Spanish AMR and UMR are able to capture for Spanish. We compare UMR with the Wein et al. (2022a) extension of AMR, which develops a corpus of approximately 500 sentences and guidelines for representing key linguistic features of Spanish in AMR. As depicted in Table 2, we find that most language-specific considerations in Spanish AMR are also included in UMR.

Verb Senses. Spanish AMR uses AnCora² verb senses, supplemented with specific senses which are not captured in the lexicon. Language-specific verb senses are used for UMR. In §5, we discuss the reliance on lexicons of both UMR and AMR.

Modality. Spanish AMR adds additional sense for *deber* (should) and *poder* (could) to mark modality. UMR marks modality through the sentence-level :MODSTR role.

Number for Persons. Spanish AMR opts against specifying number, while UMR has an additional modifying role for number of people/entities (:ref-number).

Pronoun Drop. Spanish AMR adds additional information for dropped pronouns by incorporating a sinnombre (“nameless”) concept into the graph, e.g. first-person-sing-sinnombre for implicit entities. For example, the following AMR represents the Spanish sentence *Necesito irme* (“I need to leave”), with the first-person pronoun “yo” dropped.

(1) *Necesito irme.* ‘I need to leave.’

AMR:

(n / necesitar-01
:ARG0 (f / first-person-sing-sinnombre)
:ARG1 (i / ir-05
:ARG1 f)))

UMR:

(s2n / necesitar-01
:ARG0 (s2p / person
:ref-person 1st
:ref-number Singular
:ARG1 (s2i / ir-05
:ARG1 s2p
:ASPECT Performance
:MODPRED s2n)
:ASPECT State
:MODSTR PrtAff))

UMR handles all pronouns—explicit, indexed, dropped, or implicit—via a generic concept (e.g., (p / person)) modified by :ref-person and :ref-number. There is no specific marking to indicate which of these methods of expression were used, however.

Politeness. Spanish AMR addresses politeness by adding a role relation for second person addressee. UMR adds an attribute role :polite which follows the same pattern, as follows:

(2) *usted* ‘you.FORM’

AMR:

(u / usted
:mod-polite +)

UMR:

(s3p / person
:refer-person 2nd
:refer-number Singular
:mod-polite +)

Affixes. Spanish AMR represents derivational suffixes as modifier concepts, and clitics are also treated as separate concepts.

How UMR handles derivational affixes depends on the type of affix and the annotation stage a language is undergoing. Languages undergoing stage 0 annotation (where there is no existing valency lexicon resource) may use an entire surface form (stem + affixes) as a graph predicate, or they may choose to systematically drop certain affixes as part of the lexicon-building process. Because the spirit of UMR (inherited from AMR) is to abstract away from syntactic manner of expression, lexical category-changing derivational affixes will likely be dropped from graph predicates by stage 1 annotation, with predicates coming from unified (part of speech-ambivalent) rolesets that will at that point have been created. Many other derivational affixes can now be dealt with through UMR graph structures (e.g., resemble-91 for similitive affixes). But some will need to be resolved on a language-by-language basis as part of roleset development, as occurs in cross-lingual UMR.

UMR represents inflectional affixes via :ASPECT and :MODSTR attribute roles in the sentence-level annotation and the temporal and modal dependencies at the document level (as in figure 1). The affixes themselves may also be dropped from the graph predicate as deemed appropriate for a given stage of annotation for a language.

Examples of how AMR and UMR handle derivational suffixes and clitics can be seen in (3) and (4), respectively. In (3), the diminutive suffix /-ita/ is dropped from the head concept in the graph and represented via a :mod role in both Spanish AMR and UMR. Note that UMR doesn’t have an abstract concept dedicated solely to the diminutive, and so the contents of the :mod relation will be unique to a given language, in whatever form the language deems most appropriate. The key is that the overall graph structure is the same cross-lingually.

²http://clic.ub.edu/corpus/en/ancoraverb_es

- (3) *chiquita* 'little girl'

AMR/UMR:

(c / chica
:mod (p / pequeña))

- (4) *mandarlo* 'send it'

AMR:

(m / mandar
:ARG1 (l / lo))

UMR:

(s4m / mandar :mode imperative
:ARG0 (s4p / person
:refer-person 2nd
:refer-number Singular)
:ARG1 (s4t / thing
:refer-person 3rd
:refer-number Singular)
:ASPECT Performance
:MODSTR PrtAff)

Double Negation. Double negation in Spanish can sometimes be used for emphasis, e.g. *No le dijo nada a nadie* ("She didn't say anything to anyone"). Spanish AMR specifies that double negation is treated the same as single negation (:polarity -). UMR guidelines do not state whether double negation receives special treatment, but one idea is to modify the polarity with :degree INTENSIFIER.

"Se" Usage. *Se* takes on many uses in Spanish. For AMR, there are three uses of note.

First, *se* can be used as a reflexive pronoun, annotated via reentrancy in English/Spanish AMR and UMR. For example, in 5, the reflexive verb *mirarse* (look at oneself) forces a reentrancy for *se* in both the Spanish AMR and the UMR.

- (5) *él se miraba en el espejo* 'he looked at himself in the mirror'

AMR:

(m / mirar-01
:ARG0 (e / él)
:ARG1 e
:location (s / espejo))

UMR:

(s5m / mirar-01
:ARG0 (s5e / él)
:ARG1 s5e
:location (s5e2 / espejo)
:Aspect Activity
:MODSTR FullAff)

Second, *se* can reflect a passive marker / an omitted concept (e.g. *se vende*, for sale). In this case, Spanish AMR uses the token *se* as the argument role label. UMR would annotate these passive markers as appropriate for the language and has guidelines specifically for passives. Third, *se* can be used as an impersonal pronoun (e.g. *no se debe fumar*, one should not smoke). Given that *se* is a pronoun, the second and third uses of *se* are handled in UMR using the :ref-persons concept.

Document-level representation, Scope, and Aspect. UMR expands AMR by adding annotation guidelines for document-level representation, scope, and aspect, while Spanish AMR has none of the three.

4.2 Encoding Specific Linguistic Features for Other Languages

For languages which have less syntactic similarity to English than Spanish does, some language-specific features that could be accommodated by a custom monolingual AMR-adaptation may be more straightforward to handle in UMR than others. For example, numeral noun classifiers in Vietnamese are easily covered in UMR with the numeral lattice. In Korean, UMR's flexibility towards representing affixes as concepts allows handling of case-stacking. On the other hand, specifics such as reduplicatives (in Mandarin Chinese) are not currently considered in UMR. Reduplication can occur in Mandarin by repeating a lexical unit, and can be indicative of either tentative aspects of emphasized meaning (Chen et al., 1992).

5 Challenges for UMR & Cross-lingual AMR

UMR & Cross-lingual AMR face a number of challenges when adapting to various languages, most notably in the representation of idiomatic phrases. Reliance on underlying lexicons leads to graph structural inconsistencies for parallel sentences.

Idiomatic Phrases. Idiomatic phrases are a challenge for cross-lingual AMR/UMR because of the relationship between a phrase's individual tokens and its overall meaning (Urešová et al., 2014; van der Plas et al., 2010; De Clercq et al., 2012; Kara et al., 2020). Even within a single language, it can be difficult for annotators to determine the best way to incorporate predicate argument structures associated with the specific combination of individual tokens (literal expression) and the argument

Feature	Spanish AMR	UMR
In-language verb senses	✓	✓
Modality	✓	✓
Grammatical Number	XOpted for Simplicity	✓
Pronoun Drop	✓	Not specified
Politeness	✓	✓
Affixes (Third person clitic pronouns, Suffixes)	✓	✓
Double Negation	✓Same as single negation	Not specified
Document-level representation	X	✓
Scope	X	✓
Aspect	X	✓
<i>Se</i> Usage	✓	✓Impersonal pronoun, ✓Reflexive pronouns, ✓Passive Voice

Table 2: A selection of linguistic features relevant for capturing meaning in Spanish, showing whether they are accounted for in each of the two schemata (Spanish AMR and UMR). The specific ways in which these features are accounted for in Spanish AMR and UMR are detailed in §4.1.

structure associated with the overall (idiomatic) semantics, especially when the expression is not fully compositional. Graph structures stemming from the relationships between individual tokens are, to some extent, unavoidable, and since idiomatic expressions of the same meaning can vary greatly across languages, the graph structures associated with a single meaning can also vary. An effectively cross-lingual meaning representation needs built-in considerations for addressing this challenge as uniformly as possible during annotation and parsing.

UMR has not yet established final guidelines for uniform treatment of all idiomatic phrases (but see Bonn et al. [in press](#) for further discussion), particularly during stage 0 annotation when there are no existing lexical resources to rely on that might provide a single predicate argument structure for an expression. In addition to the difficulties posed for parallel semantic representations across languages, this can also lead to inconsistencies across annotators. Still, inter-annotator agreement for small UMR annotation studies on Kukama and Arapaho, as measured by Smatch, ranges from 0.76 to 0.92, which is similar to typical AMR inter-annotator agreement scores ([Van Gysel et al., 2021b](#)).

Given that Stage 0 UMR permits annotation of tokens into multiple concepts (e.g. compound words) or of multiple tokens into a single concepts (e.g. multi-word concepts), we expect that an altered version of Smatch ([Cai and Knight, 2013](#)) will need to be adapted in order to successfully identify parallelism in meaning when quantitatively

comparing UMRs in different languages.

Reliance on English Concepts. Prior work has explored cross-lingual differences in parallel AMRs and to what extent AMR is an interlingua ([Xue et al., 2014](#); [Wein and Schneider, 2021](#)), and suggests that the AMR annotation schema may be more compatible with certain languages than others (i.e. more compatible with Chinese than Czech).

Current cross-lingual adaptations of AMR highlight this, because some cross-lingual guidelines require more changes to handle linguistic variation than others, though the structure of arguments and concepts remain largely unchanged. The approaches which use English abstract rolesets for the cross-lingual annotation (for example, *accompany-01* as the reification for the :accompanier role) exhibit significant English bias because the arguments for concepts are determined by their English usage.

AMR adaptations vary in degree of reliance on English annotations and resources, ranging from simply working with the English AMR guidelines as a baseline and extending them, to using English PropBank for sense annotation ([Miguel-Abrera et al., 2018](#)) or aligning English and Portuguese sentences and translating English annotations to their cross-lingual framesets ([Sanches Duran and Aluísio, 2015](#)). A factor that has enabled cross-lingual AMR extensions for individual languages is the existence of lexicons in those languages, such as PropBanks. This is an obstacle to AMR annotation for low-resource languages. Because many meaning representations require additional resources to

produce annotations, the lack of *prior* non-English resource work poses an issue for *future* non-English resource work (Hovy and Prabhume, 2021). This issue has been handled by UMR by developing a “road map” for annotation of low-resource languages (Van Gysel et al., 2021a).

Reliance on Frame Files. The quality/extent of the lexicon of rolesets available for a given language impacts AMR/UMR annotation. For example, Spanish AMR (Wein et al., 2022a) makes use of AnCora (Taulé et al., 2008), but despite being the most comprehensive publicly available lexical resource for Spanish, it is limited in the senses it contains, so other adaptations of AMR for Spanish have opted against its use (Migueles-Abraira et al., 2018). Thus, even with the “road map” for annotation of low-resource languages in UMR, there are complexities caused by reliance on external resources that affect UMR/AMR annotation.

Spanish AMR was forced to add a supplementary database of frame files / senses when using AnCora, and Stage 1 UMR annotation will likely also need to provide additional resources when relying on external lexicons. The UMR Writer (Zhao et al., 2021) is designed to allow annotators to add lexical entries to the roleset lexicon file used for annotation as need arises during annotation, pairing the lexicon-development process with UMR annotation. Roleset development can be incredibly complicated, however—particularly for polysynthetic and agglutinating languages like Arapaho—so this feature of the UMR-writer is a vital first step out of many when it comes to establishing a robust lexical resource.

6 Conclusion

Cross-lingual adaptations of AMR use the English annotation guidelines as a baseline, and then make a set of adaptations for linguistic features specific to the other language. The linguistic phenomena incorporated into each cross-linguistic adaptation also varies by language (as described in §2), because these phenomena are language-specific.

We conclude that UMR successfully handles the vast majority of even the more language-specific features of cross-lingual adaptations of AMR. The challenges for UMR annotation in need of further investigation and consideration include the development of quantitative metrics, which will need to account for UMR’s flexibility in multiword/affix annotation, and the complexities associated with

the generation of roleset lexicons for low-resource languages. Future work providing general insight into the morphosyntactic strategies of AMR and UMR might provide additional insight into their cross-lingual applicability.

Acknowledgements

This work is supported by a Clare Boothe Luce Scholarship and by grants from the CNS Division of National Science Foundation (Awards no: NSF_2213805, NSF_IIS 1764048 RI) entitled “Building a Broad Infrastructure for Uniform Meaning Representations” and “Developing a Uniform Meaning Representation for Natural Language Processing”, respectively. We thank anonymous reviewers as well as Andrew Cowell, Bill Croft, Jan Hajič, Alexis Palmer, Martha Palmer, James Pustejovsky, and Nianwen Xue for their helpful feedback. All views expressed in this paper are those of the authors and do not necessarily represent the view of the National Science Foundation.

References

- Rafael Anchieta and Thiago Pardo. 2018. [Towards AMR-BR: A SemBank for Brazilian Portuguese language](#). In *Proceedings of the Eleventh International Conference on Language Resources and Evaluation (LREC 2018)*, Miyazaki, Japan. European Language Resources Association (ELRA).
- Zahra Azin and Gülşen Eryiğit. 2019. [Towards Turkish Abstract Meaning Representation](#). In *Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics: Student Research Workshop*, pages 43–47, Florence, Italy. Association for Computational Linguistics.
- Laura Banarescu, Claire Bonial, Shu Cai, Madalina Georgescu, Kira Griffitt, Ulf Hermjakob, Kevin Knight, Philipp Koehn, Martha Palmer, and Nathan Schneider. 2013. [Abstract Meaning Representation for sembanking](#). In *Proceedings of the 7th Linguistic Annotation Workshop and Interoperability with Discourse*, pages 178–186, Sofia, Bulgaria. Association for Computational Linguistics.
- Julia Bonn, Andrew Cowell, Jan Hajič, Alexis Palmer, Martha Palmer, James Pustejovsky, Haibo Sun, Zdenka Urešová, Shira Wein, Nianwen Xue, and Jin Zhao. in press. Umr annotation of multiword expressions. In *Proceedings of the Fourth International Workshop on Designing Meaning Representations*, Nancy, France. Association for Computational Linguistics.
- Julia Bonn, Skatje Myers, Jens E. L. Van Gysel, Lukas Denk, Meagan Vigus, Jin Zhao, Andrew Cowell,

- William Croft, Jan Hajič, James H. Martin, Alexis Palmer, Martha Palmer, James Pustejovsky, Zdenka Urešová, Rosa Vallejos, and Nianwen Xue. 2023. [Mapping AMR to UMR: Resources for adapting existing corpora for cross-lingual compatibility](#). In *Proceedings of the 21st International Workshop on Treebanks and Linguistic Theories (TLT, GURT/SyntaxFest 2023)*, pages 74–95, Washington, D.C. Association for Computational Linguistics.
- Shu Cai and Kevin Knight. 2013. [Smatch: an evaluation metric for semantic feature structures](#). In *Proceedings of the 51st Annual Meeting of the Association for Computational Linguistics (Volume 2: Short Papers)*, pages 748–752, Sofia, Bulgaria. Association for Computational Linguistics.
- Daniel Chen, Martha Palmer, and Meagan Vigus. 2021. [AutoAspect: Automatic annotation of tense and aspect for uniform meaning representations](#). In *Proceedings of the Joint 15th Linguistic Annotation Workshop (LAW) and 3rd Designing Meaning Representations (DMR) Workshop*, pages 36–45, Punta Cana, Dominican Republic. Association for Computational Linguistics.
- Feng-yi Chen, Ruo-ping Mo, Chu-Ren Huang, and Keh-Jiann Chen. 1992. [Reduplication in Mandarin Chinese: Their formation rules, syntactic behavior and ICG representation](#). In *Proceedings of Rocling V Computational Linguistics Conference V*, pages 217–233, Taipei, Taiwan. The Association for Computational Linguistics and Chinese Language Processing (ACLCLP).
- Hyonsu Choe, Jiyoung Han, Hyejin Park, and Hansaem Kim. 2019. [Copula and case-stacking annotations for Korean AMR](#). In *Proceedings of the First International Workshop on Designing Meaning Representations*, pages 128–135, Florence, Italy. Association for Computational Linguistics.
- Hyonsu Choe, Jiyoung Han, Hyejin Park, Tae Hwan Oh, and Hansaem Kim. 2020. [Building Korean Abstract Meaning Representation corpus](#). In *Proceedings of the Second International Workshop on Designing Meaning Representations*, pages 21–29, Barcelona Spain (online). Association for Computational Linguistics.
- Orphée De Clercq, Veronique Hoste, and Paola Monachesi. 2012. [Evaluating automatic cross-domain Dutch semantic role annotation](#). In *Proceedings of the Eighth International Conference on Language Resources and Evaluation (LREC’12)*, pages 88–93, Istanbul, Turkey. European Language Resources Association (ELRA).
- Andrea Di Fabio, Simone Conia, and Roberto Navigli. 2019. [VerbAtlas: a novel large-scale verbal semantic resource and its application to semantic role labeling](#). In *Proceedings of the 2019 Conference on Empirical Methods in Natural Language Processing and the 9th International Joint Conference on Natural Language Processing (EMNLP-IJCNLP)*, pages 627–637, Hong Kong, China. Association for Computational Linguistics.
- Magali Sanches Duran and Sandra Maria Aluísio. 2011. [Propbank-br: a Brazilian Portuguese corpus annotated with semantic role labels](#). In *Proceedings of the 8th Brazilian Symposium in Information and Human Language Technology*.
- Johannes Heinecke and Anastasia Shimorina. 2022. [Multilingual Abstract Meaning Representation for celtic languages](#). In *Proceedings of the 4th Celtic Language Technology Workshop within LREC2022*, pages 1–6, Marseille, France. European Language Resources Association.
- Dirk Hovy and Shrimai Prabhumoye. 2021. Five sources of bias in natural language processing. *Language and Linguistics Compass*, 15(8):e12432.
- Neslihan Kara, Deniz Baran Aslan, Büşra Marşan, Özge Bakay, Koray Ak, and Olcay Taner Yıldız. 2020. [TRopBank: Turkish PropBank v2.0](#). In *Proceedings of the Twelfth Language Resources and Evaluation Conference*, pages 2763–2772, Marseille, France. European Language Resources Association.
- Paul Kingsbury and Martha Palmer. 2002. [From TreeBank to PropBank](#). In *Proceedings of the Third International Conference on Language Resources and Evaluation (LREC’02)*, Las Palmas, Canary Islands - Spain. European Language Resources Association (ELRA).
- Kenneth Lai, Richard Brutti, Lucia Donatelli, and James Pustejovsky. 2021. Situated umr for multimodal interactions. In *Proceedings of the 25th Workshop on the Semantics and Pragmatics of Dialogue*.
- Bin Li, Yuan Wen, Weiguang Qu, Lijun Bu, and Nianwen Xue. 2016. [Annotating the little prince with Chinese AMRs](#). In *Proceedings of the 10th Linguistic Annotation Workshop held in conjunction with ACL 2016 (LAW-X 2016)*, pages 7–15, Berlin, Germany. Association for Computational Linguistics.
- Ha Linh and Huyen Nguyen. 2019. [A case study on meaning representation for Vietnamese](#). In *Proceedings of the First International Workshop on Designing Meaning Representations*, pages 148–153, Florence, Italy. Association for Computational Linguistics.
- Abelardo Carlos Martínez Lorenzo, Marco Maru, and Roberto Navigli. 2022. [Fully-Semantic Parsing and Generation: the BabelNet Meaning Representation](#). In *Proceedings of the 60th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 1727–1741, Dublin, Ireland. Association for Computational Linguistics.
- Noelia Migueles-Abraira, Rodrigo Agerri, and Arantza Diaz de Ilarraza. 2018. [Annotating Abstract Meaning Representations for Spanish](#). In *Proceedings of the Eleventh International Conference on Language Resources and Evaluation (LREC 2018)*, Miyazaki,

- Japan. European Language Resources Association (ELRA).
- Roberto Navigli, Rexhina Billoshmi, and Abelardo Carlos Martinez Lorenzo. 2022. Babelnet meaning representation: A fully semantic formalism to overcome language barriers. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 36, pages 12274–12279.
- Roberto Navigli and Simone Paolo Ponzetto. 2010. **BabelNet: Building a very large multilingual semantic network**. In *Proceedings of the 48th Annual Meeting of the Association for Computational Linguistics*, pages 216–225, Uppsala, Sweden. Association for Computational Linguistics.
- Thi Minh Huyen Nguyen, Laurent Romary, Mathias Rossignol, and Xuân Luong Vũ. 2006. A lexicon for vietnamese language processing. *Language Resources and Evaluation*, 40(3):291–309.
- Elif Oral, Ali Acar, and Gülşen Eryiğit. 2022. Abstract meaning representation of turkish. *Natural Language Engineering*, pages 1–30.
- Martha Palmer, Paul Kingsbury, and Daniel Gildea. 2005. **The proposition bank: An annotated corpus of semantic roles**. *Computational Linguistics*, 31:71–106.
- Sameer Pradhan, Julia Bonn, Skatje Myers, Kathryn Conger, Tim O’gorman, James Gung, Kristin Wrightbettner, and Martha Palmer. 2022. **PropBank comes of Age—Larger, smarter, and more diverse**. In *Proceedings of the 11th Joint Conference on Lexical and Computational Semantics*, pages 278–288, Seattle, Washington. Association for Computational Linguistics.
- Magali Sanches Duran and Sandra Aluísio. 2015. **Automatic generation of a lexical resource to support semantic role labeling in Portuguese**. In *Proceedings of the Fourth Joint Conference on Lexical and Computational Semantics*, pages 216–221, Denver, Colorado. Association for Computational Linguistics.
- Marco Antonio Sobrevilla Cabezudo and Thiago Pardo. 2019. **Towards a general Abstract Meaning Representation corpus for Brazilian Portuguese**. In *Proceedings of the 13th Linguistic Annotation Workshop*, pages 236–244, Florence, Italy. Association for Computational Linguistics.
- Elior Sulem, Omri Abend, and Ari Rappoport. 2015. **Conceptual annotations preserve structure across translations: A French-English case study**. In *Proceedings of the 1st Workshop on Semantics-Driven Statistical Machine Translation (S2MT 2015)*, pages 11–22, Beijing, China. Association for Computational Linguistics.
- Reza Takhshid, Razieh Shojaei, Zahra Azin, and Mohammad Bahrani. 2022. **Persian Abstract Meaning Representation**. *arXiv preprint arXiv:2205.07712*.
- Mariona Taulé, M. Antònia Martí, and Marta Recasens. 2008. **AnCorà: Multilevel annotated corpora for Catalan and Spanish**. In *Proceedings of the Sixth International Conference on Language Resources and Evaluation (LREC’08)*, Marrakech, Morocco. European Language Resources Association (ELRA).
- Zdeňka Urešová, Jan Hajič, and Ondřej Bojar. 2014. **Comparing Czech and English AMRs**. In *Proceedings of Workshop on Lexical and Grammatical Resources for Language Processing*, pages 55–64, Dublin, Ireland. Association for Computational Linguistics and Dublin City University.
- Lonneke van der Plas, Tanja Samardžić, and Paola Merlo. 2010. **Cross-lingual validity of PropBank in the manual annotation of French**. In *Proceedings of the Fourth Linguistic Annotation Workshop*, pages 113–117, Uppsala, Sweden. Association for Computational Linguistics.
- Jens E. Van Gysel, Meagan Vigus, Jayeol Chun, Kenneth Lai, Sarah Moeller, Jiarui Yao, Tim O’Gorman, Andrew Cowell, William Croft, Chu-Ren Huang, Jan Hajič, James H. Martin, Stephan Oepen, Martha Palmer, James Pustejovsky, Rosa Vallejos, and Nianwen Xue. 2021a. **Designing a uniform meaning representation for natural language processing**. *KI - Künstliche Intelligenz*.
- Jens E. L. Van Gysel, Meagan Vigus, Lukas Denk, Andrew Cowell, Rosa Vallejos, Tim O’Gorman, and William Croft. 2021b. **Theoretical and practical issues in the semantic annotation of four indigenous languages**. In *Proceedings of the Joint 15th Linguistic Annotation Workshop (LAW) and 3rd Designing Meaning Representations (DMR) Workshop*, pages 12–22, Punta Cana, Dominican Republic. Association for Computational Linguistics.
- Shira Wein, Lucia Donatelli, Ethan Ricker, Calvin Engstrom, Alex Nelson, Leonie Harter, and Nathan Schneider. 2022a. **Spanish Abstract Meaning Representation: Annotation of a general corpus**. In *Northern European Journal of Language Technology, Volume 8*, Copenhagen, Denmark. Northern European Association of Language Technology.
- Shira Wein, Wai Ching Leung, Yifu Mu, and Nathan Schneider. 2022b. **Effect of source language on AMR structure**. In *Proceedings of the 16th Linguistic Annotation Workshop (LAW-XVI) within LREC2022*, pages 97–102, Marseille, France. European Language Resources Association.
- Shira Wein and Nathan Schneider. 2021. **Classifying divergences in cross-lingual AMR pairs**. In *Proceedings of the Joint 15th Linguistic Annotation Workshop (LAW) and 3rd Designing Meaning Representations (DMR) Workshop*, pages 56–65, Punta Cana, Dominican Republic. Association for Computational Linguistics.
- Nianwen Xue, Ondřej Bojar, Jan Hajič, Martha Palmer, Zdeňka Urešová, and Xiuhong Zhang. 2014. **Not an**

interlingua, but close: Comparison of English AMRs to Chinese and Czech. In *Proceedings of the Ninth International Conference on Language Resources and Evaluation (LREC'14)*, pages 1765–1772, Reykjavik, Iceland. European Language Resources Association (ELRA).

Jin Zhao, Nianwen Xue, Jens Van Gysel, and Jinho D. Choi. 2021. [UMR-writer: A web application for annotating uniform meaning representations](#). In *Proceedings of the 2021 Conference on Empirical Methods in Natural Language Processing: System Demonstrations*, pages 160–167, Online and Punta Cana, Dominican Republic. Association for Computational Linguistics.

Yuping Zhou and Nianwen Xue. 2015. The chinese discourse treebank: a chinese corpus annotated with discourse relations. *Language Resources and Evaluation*, 49:397–431.