

Linked Latent Theta Roles: A Model to Support Social Scientists with Open-ended Exploration of Framing

Anonymous ACL submission

Abstract

Computational research has developed techniques to classify frames in text. However, these techniques may be less useful for supporting researchers in *exploratory analysis of framing* as an act of meaning construction. To address this gap, we introduce Latent Linked Theta Roles (LLTR), a model based on linguistic attributes relevant to framing language. Rather than identifying frames *per se*, the LLTR model highlights linguistic patterns that might be indicative of framing, thus supporting researchers in conducting open-ended, exploration of framing. A qualitative human-subject study compares this novel model against two baseline models, demonstrating that LLTR is more effective in assisting researchers with this exploratory task.

1 Introduction

“Facts have no intrinsic meaning. They take on their meaning by being embedded in a frame or story line that organizes them and gives them coherence.” (Gamson and Modigliani, 1989, p.157). The processes involved in these meaning constructions are referred to, by sociological researchers, as *framing* (Gamson, 1989; Scheufele, 1999; Benford and Snow, 2000; Druckman, 2001), and they are evidenced in linguistic patterns. For instance, the phrases “the soldier shot a bystander” and “a bystander was shot by the soldier” denote the same information, but they frame the situation differently in terms of responsibility.

Computational research has focused on analyzing linguistic techniques to examine framing (Card et al., 2016; Baumer et al., 2015; Walter and Ophir, 2019; Naderi and Hirst, 2017). Most such work offers computational models to identify frames as discrete, distinct entities that could be present or absent in a corpus.

However, less computational work has investigated models specifically designed to support social science researchers in *exploratory analysis*

of framing. The inherently subjective nature of framing analysis (Schön and Rein, 1994; Kuypers, 2010; Van Gorp et al., 2010) makes prior computational work on directly labeling frames (cited above) poorly suited to this task.

Instead, this paper describes the design, implementation, and evaluation of a novel computational model that captures linguistic evidence indicative of framing. The model does not identify framing *per se*, but rather can assist researchers in conducting exploratory analysis of framing. The design of the model draws on insights from the definition of framing in sociological research (Gamson, 1989; Scheufele, 1999; Benford and Snow, 2000; Druckman, 2001), and from prior computational work focused on *identifying frames* (Baumer et al., 2015; Card et al., 2016). Specifically, this approach simultaneously models two linked distributions per topic, one for the grammatical relations in which a topic’s terms occur, and another one for the words that co-occur with those topic terms within these grammatical relationships. Taking inspiration from LinkLDA (Erosheva et al., 2004) and prior extensions thereof (Ritter et al., 2010), we refer to this as the *Linked Latent Theta Role (LLTR)* model.

The paper evaluates this novel model by examining its utility in helping guide researchers’ attention to language indicative of framing. Our LLTR model is compared against two simpler baseline models: standard LDA (latent Dirichlet allocation) (Blei et al., 2003), which has previously been applied to analyzing framing (Walter and Ophir, 2019), and LDA that accounts for grammatical relations by simply appending to each word token the grammatical relation in which the word occurs (e.g., direct object of a verb).

Given the importance of “relevant human readers” in model assessments (Hoyle et al., 2021), we compared these three models via a human subjects evaluation with researchers who have familiarity with framing. Furthermore, given the subjective

nature of framing analysis (Shön and Rein, 1994; Kuypers, 2010; Van Gorp et al., 2010), we leveraged qualitative methods to understand the criteria by which participants assessed each model, as well as which models they found preferable.

Evaluation results indicate that participants perceived LLTR as the most effective model for exploratory framing analysis. They noted LLTR provides broader corpus overview, and readily offers framing evidence across documents. Participants reported greater confidence in their framing analysis using LLTR, compared to the other two base models. Despite an initial learning curve, participants reported that once familiar with LLTR’s components, it facilitated a smoother, more effective process for finding framing evidence, compared against two simpler models.

Thus, this paper both posits, and demonstrates the viability of an alternative approach to computational techniques to support framing research. That is, rather than trying to identify frames for researchers (e.g., Card et al., 2016; Naderi and Hirst, 2017; Morstatter et al., 2018), it argues and demonstrates computational techniques can instead draw attention to linguistic patterns potentially indicative of framing, thereby assist human researchers to interpret framing.

2 Related Work

2.1 The concept of framing

The concept of framing is studied in different fields (Scheufele, 2000; Tversky and Kahneman, 1985; Goffman, 1974; Kahneman and Tversky, 1984). This work adopts the definition of framing from sociological studies (e.g., Gamson, 1989; Benford and Snow, 2000). These studies define framing as a set of processes by which people come to interpret and understand world’s events. Framing performs many functions, including what counts issues, how the causes are diagnosed, moral judgments being made about events under discussion, and what remedies are suggested (Gamson, 1989; Entman, 1993).

Sociological researchers argue that frames are not fixed or pre-categorized static, discrete entities. They rather focus on how framing, evident in language, operate within the dynamic interpretation and construction of meaning (Gamson and Modigliani, 1989; Benford and Snow, 2000), often using exploratory, open-ended methods.

2.2 Prior Computational Analysis of Framing

Prior work examined a variety of techniques to analyze framing (Baden, 2018; Van Atteveldt, 2008; Card et al., 2016; Sturdza et al., 2018; Ziems and Yang, 2021; Khanehzar et al., 2019; Mendelsohn et al., 2021; Morstatter et al., 2018). This paper focuses on unsupervised topic modeling, both for its popularity in framing research (e.g., Walter and Ophir, 2019; Jacobi et al., 2018; Ylä-Anttila et al., 2022; DiMaggio et al., 2013), and for its alignment with the linguistic attributes relevant to the language of framing (discussed in §3).

Despite this popularity, a standardized approach for utilizing topic modeling to investigate framing has not been established. For example, (DiMaggio et al., 2013, p.578) state that “many topics may be viewed as frames.” Ylä-Anttila et al. (2022), on the other hand, discuss the use of “topics” as a proxy for “frames” is conditioned on three criteria: defining framing as a connection between concepts, a subject-specific corpus, and validation against existing frame analyses. Both these studies suggest a direct link between frames and topics. However, topics are unintelligible linguistic patterns of word co-occurrence (Blei et al., 2003). Thus, such strict mapping is potentially reductive, obscuring framing complexity. Specifically, while topics can identify discussed issues, they often lack nuanced information about causality, interpretation, or suggested remedies (Ali and Hassan, 2022)

Only recently some researchers have moved beyond mapping topics to frames. For example, Walter and Ophir (2019) argue that frames can be considered as *communities/clusters* in a networks of topics. While effective for capturing established frames, their approach offers limited insights into framing packages (Ali and Hassan, 2022). Card et al. (2016) argue that understanding framing requires attention to *narratives*, particularly the entities involved. They contend that relying solely on word co-occurrence patterns, i.e., topics, is insufficient for identifying narratives. Instead, they proposed an unsupervised model that clusters characterizations of entities into personas¹. This approach produces interpretable clusters (i.e., persona), that effectively predict a set of predefined frames.

In summary, the reviewed computational techniques to analyze framing are primarily classifica-

¹The concept of persona is introduced by Bamman et al. (2013). Unlike Bamman et al. (2013), though, Card et al. (2016) allow personas to account for entities other than predefined characters, such as institutions, objects, and concepts.

tory in nature, aiming to directly identify frames. These approaches, however, are less appropriate when supporting researchers in *exploratory* analysis of framing. Furthermore, these approaches mostly focus on studying only *words* as framing evidence. However, there could exist other linguistic attributes that might provide insights not only *what* issues are discussed, but also other functions by which framing performs, such as *how* issues are discussed, what arguments are conveyed.

3 Linguistic Attributes Relevant to Framing Language

This section details linguistic attributes relevant to framing, which inform models design in §4.

3.1 Word Choice

Framing literature suggests framing often manifest through particular “keywords”, and “stock phrases.” (Entman, 1993), as well as “catchphrases” and “exemplars” (Gamson and Modigliani, 1989).

Indeed, the definition of framing highlights the importance of word choice as well. Specifically, word choice can help infer the events under discussion, the issues highlighted around those events, and the potentially responsible parties involved. Furthermore, the word choice can provide insights around how events and their associated issues are *labeled*. Labeling is indeed an important component of framing (Lau and Schlesinger, 2005). For example, in the case of the COVID-19 vaccines, word choice can signal if vaccines are labeled as a societal right (associated with words such as fundamental rights, societal rights), or as a marketable commodity (associated with words such as consumer choice, private insurance).

Prior computational framing research often uses word-based features (Baumer et al., 2015; Naderi and Hirst, 2017; Morstatter et al., 2018). Comparing different features, Baumer et al. (2015) found that lexical features (unigrams, bigrams, trigrams) were important indicators of framing language.

3.2 Latent Themes (i.e., topics)

Word co-occurrence patterns in a corpus can reveal framing. These patterns are often analyzed for latent themes using topic modeling (Blei, 2012; Roberts et al., 2014; Lucas et al., 2015).

Prior work utilized latent themes (i.e., topics) to identify dominant frames (discussed in §2.2). For example, Walter and Ophir (2019) suggest that

latent topics helps identify frame devices, including word choices, metaphors, or catchphrases.

Without making any restrict connection between topics and framing, this paper posits that examining topics in a corpus might provide *evidence* that can help attend to interpretive packages (i.e., frames). Specifically, instead of labeling latent topics as in the work presented by Walter and Ophir (2019), this work utilizes these topics for exploratory analysis of framing, as outlined in §2.1.

3.3 Grammatical Relationships

While knowing which groups of words co-occur can be informative, framing may also be indicated by the relationships among those words. The grammatical structure of sentences may help indicate those relationships (Pan and Kosicki, 1993; Hallahan, 1999; Fairclough, 2013).

Indeed, few prior computational work demonstrates that grammatical structures are important indicators of frame evoking language (Baumer et al., 2015; Recasens et al., 2013). For example, Baumer et al. (2015) show that the grammatical relations in which words appear within a document are important when inferring frames within a document.

While Baumer et al. (2015) focus on identifying frames in a classificatory approach, this paper posits that grammatical relationships may similarly be important for exploratory analysis of framing. Relevant to the perspective on framing adopted here, in addition to capturing *what* people discuss (captured via *word choice* and *latent themes*), it is important to explore *how* people discuss an event. Accounting for grammatical structures between words might address this aspect.

4 Model Designs for Framing

This section first describes and motivates the two simpler baseline models against which LLTR is compared. It then presents both the details of and the rationale behind the LLTR model.

4.1 First Baseline Model: LDA

4.1.1 Motivations

Prior work utilizes topic modeling to infer dominant frames (e.g., Walter and Ophir, 2019; Ylä-Anttila et al., 2022). However, it remains unclear whether and how topic modeling can be leveraged to explore framing as a meaning-making process (see §2.1). This paper posits that LDA provides insights into *word choice* and *underlying themes*

(see §3), thus offering valuable linguistic evidence to support exploratory analysis of framing.

4.1.2 Model Description

LDA identifies latent topics in large corpus of documents (Blei et al., 2003; Blei, 2012). Each topic (i.e., themes) is represented as latent probability distribution over all the words in the vocabulary of a corpus. LDA allows for multiple memberships of words in various topics. Thus, the same word can be interpreted differently (implicitly, by a human reader) depending on the context (i.e., the probabilities of other words in the topic) (Blei, 2012; DiMaggio et al., 2013; Walter and Ophir, 2019).

4.2 Second Baseline Model: LDA-GR

4.2.1 Motivation

As noted above (§3.3), *grammatical relationships* can be indicative of framing. However, most prior topic modeling techniques account only for *word choice* and co-occurrence patterns, i.e., *latent themes*. Our second baseline, named the Latent Dirichlet Allocation-Grammatical Relationship model, i.e., LDA-GR, uses a simple extension of LDA to account for grammatical relationships.

4.2.2 Model Description

LDA-GR replaces each word token with a concatenation of the token itself and its grammatical role. After parsing each document (Manning et al., 2014), for each word token w in document d , a tuple of $\langle w, reln.role \rangle$ is created, wherein $reln$ is the typed dependency of the word w in the document d , and $role$ specifies the role of the word w in the typed dependency of $reln$. For example, tuples for the sentence “Science defeated Covid-19” would include $\langle \text{defeated}, nsubj.gov \rangle$ and $\langle \text{science}, nsubj.dep \rangle$, among others, since “science” is the nominal subject of “defeated.” LDA-GR uses the same model structure as LDA but is trained on these $\langle w, reln.role \rangle$ tuples rather than on word tokens.

4.3 Novel Model: Linked Latent Theta Roles (LLTR)

4.3.1 Motivation

Although LDA-GR incorporates grammatical relationships, LDA-GR’s simplistic, one-to-one mapping between word tokens and grammatical relationships increases vocabulary size and sparsity, potentially reducing topics coherence (Blei and Lafferty, 2006; Popescul et al., 2013). In addition,

it does not capture which governors and dependents actually co-occur (e.g., which nominal subjects go with which verbs). Furthermore, not all social scientists have the formal linguistics background to be familiar with relations such as “ccomp” (causal complement) or “xsubj” (controlling subject).

To address the aforementioned challenges with LDA-GR, we designed and developed the linked latent theta role model (LLTR). Instead of capturing one-to-one correspondence between each word token and grammatical relationship, LLTR learns distributions over the grammatical relations in which topic terms occur. Thus, it can account for grammatical relationships without increasing the vocabulary size or, thereby, sparsity. It also captures syntactic variations that might be semantically equivalent but have connotative differences relevant to framing, e.g., “Science defeated COVID-19” vs. “COVID-19 was defeated by science.”

Much in the same way that LDA uses latent topic variables to represent probability distributions over words (Blei et al., 2003), LLTR uses latent variables, which we term *theta roles*², to capture probability distributions over the set of grammatical relationships in which topic words occur. In addition, to facilitate finding connections between topic terms (e.g., which nominal subjects co-occur with which verbs), LLTR also captures distributions over the second argument (either the governor or the dependent) with which topic terms occur within a grammatical relationship. In the above example, LDA-GR strictly enforces the relationships $\langle \text{science}, nsubj.dep \rangle$ and $\langle \text{defeated}, nsubj.gov \rangle$, but it captures no relationship between these two tokens. In contrast, LLTR uses latent theta roles to capture a probabilistic three-way relationship among “science,” the *nsubj* relation, and “defeated,” as described next.

4.3.2 Model Description

Linked theta role components include (1) distributions over grammatical relationships, and (2) distributions over associated arguments that appear in these grammatical relationships within the topic’s context. Analogous to LinkLDA (Ritter et al., 2010) (based on (Erosheva et al., 2004)), LLTR employs a linked latent variable to enable learn-

²Similar to theta roles in English (Aronow, 2016), the latent theta roles in LLTR are intuitively captured by syntactic structures. However, these variables are designed only to model probability distributions over grammatical relationships, without any direct mapping to syntactic or semantic constructs.

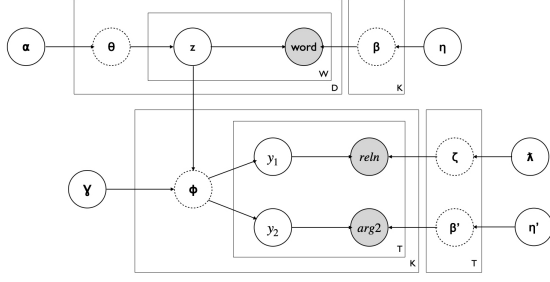


Figure 1: LLTR's plate diagram. Within each topic, linked theta roles model the probability of grammatical relations and the co-occurring word for those relations.

ing associated pairs of grammatical relationships and arguments that appear in those grammatical relationship (*reln*, *arg*). To do so, rather than requiring both components to be generated from one possible pairs of $|T|$ multinomials (ζ_t, β'_t), LLTR allows these component (i.e., the grammatical relationships and the associated arguments) to be drawn from $|T|^2$ possible pairs. However, to favor states where the grammatical relation *reln* and the arguments component *arg* are derived from related theta role assignments, this model employs a sparse prior over the theta role distributions. The terminologies and generative story of the LLTR model are defined below.

Definitions First, let there be K latent topics, where each topic β_k is a multinomial over the V words in the vocabulary (Blei et al., 2003), drawn from a Dirichlet parameterized by η (i.e., $\beta_k \sim \text{Dir}(\eta)$). For each topic, define T latent theta roles ϕ_t , where each theta role has a set of two multinomial ϕ_{1t} and ϕ_{2t} , corresponding to the two component of theta role (i.e., grammatical relationships *relns*, and argument components *arg*). Specifically, ϕ_{1t} is a multinomial distribution over the K latent topics for the first component of the theta role t , which is associated with R numbers of grammatical relationships *reln*. Each grammatical relationship is drawn from a Dirichlet distribution, parameterized by γ (i.e., $\phi_1 \sim \text{Dir}(\gamma)$). Within the ϕ_{1t} matrix, each row represents the topic distribution of the theta role t over the K latent topics. ϕ_{2t} , on the other hand, is a multinomial distribution over the K latent topics for the second component of the theta role t , which is associated with A numbers of argument components, *arg*. These arguments are drawn from a Dirichlet distribution parameterized by γ (i.e., $\phi_2 \sim \text{Dir}(\gamma)$). Within the ϕ_{2t} matrix, each row represents the topic distribution of the

theta role t over the K latent topics.

In the generative process, for each document d_i :

- Select document length $N \sim \text{Poisson}(\xi)$.
- For each word w_j to w_N :
 - Draw a topic assignment z with corresponding multinomial distribution over latent topics from the θ matrix, based on $P(z|\theta, d_i)$
 - Conditioned on the topic z , draw a theta role y_1 with corresponding distribution from the ϕ_1 matrix (i.e., $y_1 \sim \text{Multinomial}(\phi_1)$).
 - Choose the grammatical relationship *reln* from $P(\text{reln}|y_1, \zeta)$
 - Conditioned on the topic z , draw a theta role y_2 with corresponding distribution from the ϕ_2 matrix (i.e., $y_2 \sim \text{Multinomial}(\phi_2)$).
 - Choose the argument component a from $P(a|y_2, \beta')$
 - For the topic z drawn in the previous step, choose w_j from $P(w_j|z, \beta)$

The inference process The inference for the topic-word distribution, i.e., β , the model adopts the process in LDA (Blei et al., 2003), hence omitted for simplicity. For the inference on the probability distribution of theta role components, including ϕ_1 and ϕ_2 , collapsed Gibbs sampling is employed as follows (Griffiths and Steyvers, 2004; Geman and Geman, 1984). At each iteration, for each word w , provided topic z is selected, we sample theta role y_1 from the grammatical relationship component of theta role Φ_1 as follows.

$$P(y_1|\text{reln}_i, \Phi_1, z) \propto P(\text{reln}_i|y_1) * P(y_1|\Phi_1, z)$$

$$P(\text{reln}_i|y_1) = \frac{\text{Count}_{\text{reln}_i, y_1} + \lambda}{\sum_{i=1}^R \text{Count}_{\text{reln}_i, y_1} + \lambda * R}$$

$$P(y_1|\Phi_1, z) = \frac{\text{Count}_{y_1, z} + \gamma}{\sum_{j=1}^T \text{Count}_{y_1, j} + \gamma * T}$$

$\text{Count}_{\text{reln}_i, y_1}$ is the count of all words whose grammatical relationship is reln_i and the first argument of their theta role is y_1 . $\sum_{j=1}^R \text{Count}_{\text{reln}_i, y_1}$ is the same count, summing over all the R grammatical relationships *reln*. Specifically, the probability that theta role y_1 is selected for the first component of theta role based on the Φ_1 distribution, provided that topic z is selected is proportional to the probability of grammatical relationship reln_i belong to the theta role y_1 , times the contribution of theta role y_1 for the topic z .

Following a similar approach, the theta role y_2 is sampled from the second argument component of theta role Φ_2 .

$$P(y_2|arg_i, \Phi_2, z) \propto P(arg_i|y_2) * P(y_2|\Phi_2, z)$$

$$P(arg_i|y_2) = \frac{Count_{arg_i, y_2} + \beta'}{\sum_{j=1}^A Count_{arg_j, z} + \beta' * A}$$

$$P(y_2|\Phi_2, z) = \frac{Count_{y_2, z} + \gamma}{\sum_{j=1}^T Count_{y_2, z} + \gamma * T}$$

Employing the previous notation, $Count_{arg_i, y_2}$ is the count of all words that are observed in a grammatical relationship with the argument arg_i , and $\sum_{j=1}^A Count_{arg_j, z}$ is the same count, summing over all the possible A arguments within the corpus. The probability of choosing y_2 as the second component of latent theta role for the component arg_j , is proportional to the probability of the argument arg_j belongs to theta role y_2 , times the contribution of this theta role y_2 for the assigned topic z .

5 Model Evaluation: Methods

5.1 Human-subject Evaluation: Motivations

This paper uses a qualitative human-subject evaluation to assess each model for two main reasons. First, there is little evidence that metrics designed for assessing topic quality (Röder et al., 2015) will align with human perceptions of relevance to framing. Indeed, prior work has suggested that well-establish topic modeling metrics may diverge from human perceptions (Hoyle et al., 2021; Hosseiny Marani et al., 2022). Second, there does not exist a one-to-one mapping between topics and framings. That is, topics, and their associated components (e.g., grammatical relationships in LDA-GR and co-occurring terms in LLTR) are intended to be used collectively to help a researcher explore framing. Thus, even if a relevant topic-based metric existed, there is no guarantee that the aggregate of such a metric applied to individual topics would align with human perception of efficacy of a model as a whole.

Indeed, a number of other studies also employed human-subject studies to evaluate models for complex concepts (Hoyle et al., 2021; Poursabzi-Sangdeh et al., 2016; Smith et al., 2017).

5.2 Participants

To involve “relevant human readers” in the evaluation (Hoyle et al., 2021), used convenience sampling (Jager et al., 2017) to recruit ten researchers, ranging from graduate students to associate professors, experienced with analyzing framing.

5.3 Study Material

5.3.1 Dataset

All the three models are trained on a COVID-19 dataset. This subject is chosen given the particular importance of framing for health communications (Park and Reber, 2010; Salovey and Wegener, 2003; Guenther et al., 2021). This dataset contains 3,655 COVID-19 news articles collected from health department sources across 30 U.S. states where this content could be readily collected.

Initial observations indicated a potential skew in captured topics towards specific U.S. states, e.g., a topic with top terms such as *texas*, *abbott*, *austin*, *greg*, *houston*, or *lamont*, *ned*, *connecticut*. To mitigate such topics, we implemented down-sampling (Thompson and Mimno (2018)), treating each state as an author. This approach aims to enhance the representativeness of topics across various sources (i.e., states within our dataset).

5.3.2 Experiment Set Up

Coherence analysis (Newman et al., 2010; Li et al., 2024) identified six topics ($K=6$) to provide the most cohesive topics for the corpus, adopted in three models. To determine the optimal number of linked latent theta roles (T), this paper followed the approach outlined in Bamman et al. (2014), testing $T \in \{5, 10, 15, 20\}$. Selection of T was based on co-occurring words cohesion, which suggests five linked theta roles ($T=5$). For the evaluation study, the top three most coherent, and thematically similar topics were chosen per model.

Each model’s topic terms and components are illustrated with example documents showing their appearance within the topic context. For LDA, topic terms are augmented with example documents, provided that example documents have high probability for the given topic. For LDA-GR, each pair of topic term and its grammatical relationship is augmented with example documents, within the topic’s context. For LLTR, each triple of topic term, its co-occurring words, and their grammatical relationships, are augmented with example documents wherein these elements jointly manifest, within topic’s context. Given the linked theta role construct is captured at topic level, a post-hoc analysis is conducted to capture the distributions of theta role components (i.e., grammatical relationship and argument components), for each topic term, given the topic-level linked theta role construct. The constant probability of linked theta role within each

topic is omitted for brevity.

$$P(\text{reln}_i|w_j, \Phi_1) \sim P(\text{reln}_i|w_j) * P(w_j|\Phi_1)$$

$$P(\text{arg}_i|w_j, \Phi_2) \sim P(\text{arg}_i|w_j) * P(w_j|\Phi_2)$$

We also developed a visual interface to support the use of these models, as detailed in Appendix B.

5.4 Procedure

The evaluation approach consists of a two-phase study, outlined below. The study was approved by the IRB [X] at [Anonymous]. Participants were compensated with \$50.00 Amazon gift card.

5.4.1 Phase 1: Preliminary Task

The preliminary task engaged participants with model results. Following a within-subject design, each participant reviewed two models, to facilitate model comparison with fewer participants.

Upon consenting to participate, participants were provided with framing definition. Next, to ensure engagement with the models' results, for each model, participants were tasked to review the results, and answer a series of questions on framing. The questions focused on the functions by which framing performs (e.g., issues discussed, causes identified). These responses were not directly analyzed, but rather were used to scaffold succeeding follow-up, semi-structured interview.

5.4.2 Phase 2: Interview Study

Once participants completed the preliminary tasks, to gain more understanding about their experience, they were invited to a semi-structured interview.

Participants received their phase 1 responses prior to the interview, which formed the basis of discussion during the interview. They were encouraged to share how they leveraged the model's results to attend to respond to the framing questions in phase 1, such as how they inferred the *issues*, *causes*. If a participant did not respond to a question in phase 1, they were asked to describe any attempts they made, regardless of success.

Transcripts of the interviews were analyzed using thematic analysis (Braun and Clarke, 2006; Lofland et al., 2022). This qualitative method requires iteratively reviewing transcripts of interviews to identify salient themes. Here, we sought themes pertaining to how each participant determined the effectiveness of the linguistic patterns identified by each of the discussed models in assisting them to analyze framing processes.

6 Models' Evaluation: Results

This section presents semi-structured interview results (§5.4.2) assessing participants' perceived efficacy of each model for exploratory framing analysis. Specifically, given the goal of this evaluation, it focuses on how participants used model results and evaluated their efficacy and utility in facilitating framing analysis, rather than the specific framings identified. Thematic analysis (Braun and Clarke, 2006; Lofland et al., 2022) of the interviews suggests participants evaluated the models based on four criteria, described below.

6.1 Identified Evaluation Criteria

Context concerns the extent to which each model provides information about how each linguistic pattern (e.g., topic terms, or argument words) appears in its immediate sentence, within an entire document, and across the corpus more broadly.

Clarity focuses on how clearly the participants can understand the meaning of topic terms, the connections between example documents, and readily find the framing evidence in the corpus.

Confidence reflects participants' perception of each model's results' representativeness of the corpus, and their certainty about the thoroughness and evidential basis of their inferred framings.

Curve, i.e., learning curve, refers to the time and effort the participants need to spend to both learn different components of each model's results, and to then use those components to analyze framing.

6.2 Overview of Models' Evaluation

This section details participants evaluations of each models, summarized in Table 1, Appendix A.

LDA Assessment: Participants found that LDA fell short in providing sufficient *context* required for framing analysis. They mentioned the example documents "felt random" (P1³), and there was not enough "diverse contexts" (P4). Thus, participants could not infer *why* each topic term is important and *what arguments* they convey (e.g., "I know that communities is important. I don't know exactly *how* they're talking about communities.", P2).

Participants frequently found LDA's connection between topic terms and documents unclear. To *clarify* these links, they often had to read full documents. For example, P2 explained, "Just having

³Quotes are attributed using random participant IDs.

access without related words leaves me wondering, okay, what access? It could mean so much. So I had to read more.” However, this effort didn’t always clarify the ambiguity (P2, P4)."

The perceived randomness of provided context led to lower participant *confidence* about LDA results being representativeness of the corpus. For instance, P1 mentioned that “it’s hard to tell if this is indicative of bigger things in these topics just because I think there’s like one or two examples where it sticks out and, it’s much harder I think to assess like, does this cover this whole data set? [...] it felt a little bit random”. LDA’s simplicity offered the easiest learning curve. This simplify, made it less useful for finding framing evidence.

LDA-GR Assessment: Participants found LDA-GR lacked diverse *contexts*, yet grammatical relationships prompted more intentional reading. P8 noted seeing a term in different relationships emphasized its importance. However, they felt LDA-GR missed distributed framing evidence. P6 explained “you can get the framing at a sentence level, and I think this model would be useful for that, but it’s gonna miss more subtle or ambiguous frames, that aren’t necessarily explicitly stated”

LDA-GR was slightly more effective in clarifying terms connections, due to inclusion of grammatical relationships (P5, P6). However, given the effort required to interpret these relationships, several participants reported abandoning this feature (P3, P5, P8). Participants noted lack of sufficient evidence to assess whether the inferred framings were well supported, made them less *confident* in the thoroughness of their analysis. P6 hesitated relying on grammar to infer framing, stating “I’m not huge into focusing on grammar, because I study social media and people are terrible about grammar.”

LDA-GR exhibited a high learning *curve*, due to the inclusion of grammatical relationships (e.g., “as a native English speaker, I haven’t thought about grammar since secondary school.”, P6)

LLTR Assessment: Participant reported LLTR exhibited the highest efficacy in offering diverse and readily connected *contexts*. They mentioned co-occurring terms were highly effective for understanding of the broader arguments being made across the corpus (P1, P2, and P4, P6, P7) (e.g., “these different [co-occurring terms] words that go together I think provides what feels like a more holistic overview of what is in the data.”, P1).

Participants reported that LLTR made *clarifying* meaning of terms and their connections simpler than LDA and LDA-GR. This improved usability was due to LLTR’s inclusion of co-occurring terms (e.g., “I didn’t need to read very much. I could just tell from even this list of co-occurring words”, P2).

LLTR fostered increased *confidence* in representativeness of its result and participants’ certainty in their analyses. For instance, P1 and P2 mentioned that LLTR helped them understand the overarching argument and increased their confidence in inferring framing from a wider range of evidence compared to LDA and LDA-GR. Similarly, P1 emphasized that with LLTR, it is not just “volume” of examples, but also “the breadth of different documents” that “gives more confidence that you’re getting a fuller picture of what this [corpus] is”.

Participants reported LLTR had the greatest *learning curve* due to its increased number of components. However, once participants were familiar with it, LLTR sped up and smoothed the process of examining framing evidence. P1 noted “there was maybe a little bit of a learning curve to figure out how to use it. The upside of all that complexity is that there’s sort of a lot more nuance here. [...] [LLTR] topics did sort out a lot more discreetly. It was much easier to sort of see them as different things. And all of the examples I think were incredibly useful to dig in and get a better sense of what these words were doing.”

7 Conclusion and Future Work

This paper contributes LLTR, a model to support social science researchers’ exploratory analysis of framing, as well as an evaluation demonstrating the model’s efficacy. As a technical contribution, the model’s integration of word choice, themes, and grammatical relationships accounts for the varied, distributed evidence that researchers use in exploratory analysis of framing. The evaluation demonstrates how these aspects of LLTR’s model design help researchers account for broader contexts when exploring framing. It also offers a set of criteria that may be useful for future work.

Future work would benefit from designing and testing other approaches that, rather than classify frames, help draw researchers’ attention to patterns of language potentially indicative of framing. The positive results presented above illustrate the viability of this novel yet promising direction for computational approaches to framing.

8 Limitations and Ethical Considerations

This section acknowledges limitations both in the LLTR model’s design and in the evaluation study. Throughout, it notes how these limitations suggest valuable avenues for future research.

The first limitation pertains to the evaluation study. Specifically, the study was conducted using convenience sampling. This sampling technique was chosen due to the necessity of involving researcher participants, a “relevant human readers” (Hoyle et al., 2021) for the models examined in this paper, and the difficulty of recruiting a sufficient number of framing researchers to assess these models. It is thus important for future work to conduct further human-subject studies to shed light into how these models might be evaluated by other researchers.

In addition, the LLTR model’s design (similar to LDA and LDA-GR) places equal importance on different parts of an article in terms of framing language. However, some comments from our study participants suggested that framing language may be distributed unevenly across document sections (e.g., titles, introductions, bodies, conclusions). Thus, future research should explore implementing variable weighting to reflect the differing probabilities of finding framing evidence in each section, as informed by expert knowledge of researcher attention and empirical studies.

Moreover, the LLTR model, in its current state, is tested on grammatically well-structured documents (i.e., reports from the Department of Health across 30 US states). Thus, future research should assess the LLTR model’s effectiveness at identifying framing evidence in less structured data, including spoken language, slang, and short documents. Specifically, alternative grammatical parsers designed for short texts should be explored (Kong et al., 2014; Liu et al., 2018).

Furthermore, the presented study selects the LLTR’s hyperparameter, i.e., number of linked theta role number (T) using an exploratory approach. While this approach was similarly taken in other studies that introduce a new latent construct (Bamman et al., 2014), future research is encouraged to investigate the feasibility of using coherence metrics to guide the selection of T , analogous to their use in determining number of topics (i.e., K).

Also, the conducted human-subject study did not compare the efficiency of LLTR with manual

content analysis. That said, it is important for future work to compare framing analysis with and without the support of the LLTR model, and to explore what aspects of framing LLTR model might (or might not) miss compared to manual coding of content, and what aspects it might shed light into that manual coding might miss. Such exploration can provide insights into what a researcher can gain using this model, and what they might be missing when using this model to explore framing processes.

Lastly, the LLTR model presented in this paper is compared against LDA and LDA-GR (i.e., an extension of LDA that incorporates grammatical relationships). Notably, we did not compare LLTR with more recent contextualized topic modeling techniques, such as BERTopic (Grootendorst, 2022), for two primary reasons. First, several recent studies have suggested that neural topic models, including contextualized approaches, do not consistently outperform probabilistic methods such as LDA (Doogan and Buntine, 2021; Hoyle et al., 2021; Hosseiny Marani, 2025). Second, even if neural models consistently outperformed probabilistic models, current contextualized topic modeling techniques, in their standard formulation, do not explicitly model grammatical relationships during topic generation. As shown above, such relationships were central to participants’ assessments of the models during our evaluation study. That said, we acknowledge the importance of future work to directly compare and contrast the efficacy of LLTR against state-of-the-art contextualized topic models (Bianchi et al., 2020b,a). For example, future research should investigate whether attention-based neural networks, such as masked language models, can implicitly capture and leverage grammatical relationships relevant to framing without explicit grammatical knowledge beyond attention-driven word associations. Such an exploration could potentially yield novel insights into the extent to which attention mechanisms learn linguistic nuances pertinent to framing. Alternatively, future work could explore incorporating explicit grammatical information into attention-based models to potentially enhance contextualized topic modeling for tasks like framing analysis, bridging the advancements in both probabilistic and neural approaches.

References

- Mohammad Ali and Naeemul Hassan. 2022. A survey of computational framing analysis approaches. In *Proceedings of the 2022 Conference on Empirical Methods in Natural Language Processing*. 9335–9348.
- Robin Aronow. 2016. Semantics: Thematic Roles. <https://www.linguisticsnetwork.com/semantics-thematic-roles/>
- Christian Baden. 2018. Reconstructing frames from intertextual news discourse: A semantic network approach to news framing analysis. In *Doing News Framing Analysis II*. Routledge, 3–26.
- David Bamman, Brendan O’Connor, and Noah A Smith. 2013. Learning latent personas of film characters. In *Proceedings of the 51st Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*. 352–361.
- David Bamman, Ted Underwood, and Noah A Smith. 2014. A bayesian mixed effects model of literary character. In *Proceedings of the 52nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*. 370–379.
- Eric Baumer, Elisha Elovic, Ying Qin, Francesca Polletta, and Geri Gay. 2015. Testing and comparing computational approaches for identifying the language of framing in political news. In *Proceedings of the 2015 conference of the North American chapter of the Association for Computational Linguistics: human language technologies*. 1472–1482.
- Robert D Benford and David A Snow. 2000. Framing processes and social movements: An overview and assessment. *Annual review of sociology* 26, 1 (2000), 611–639.
- Federico Bianchi, Silvia Terragni, and Dirk Hovy. 2020a. Pre-training is a hot topic: Contextualized document embeddings improve topic coherence. *arXiv preprint arXiv:2004.03974* (2020).
- Federico Bianchi, Silvia Terragni, Dirk Hovy, Debora Nozza, and Elisabetta Fersini. 2020b. Cross-lingual contextualized topic models with zero-shot learning. *arXiv preprint arXiv:2004.07737* (2020).
- David Blei and John Lafferty. 2006. Correlated topic models. *Advances in neural information processing systems* 18 (2006), 147.
- David M Blei. 2012. Probabilistic topic models. *Commun. ACM* 55, 4 (2012), 77–84.
- David M Blei, Andrew Y Ng, and Michael I Jordan. 2003. Latent dirichlet allocation. *Journal of machine Learning research* 3, Jan (2003), 993–1022.
- Virginia Braun and Victoria Clarke. 2006. Using thematic analysis in psychology. *Qualitative research in psychology* 3, 2 (2006), 77–101.
- Dallas Card, Justin H Gross, Amber Boydstun, and Noah A Smith. 2016. Analyzing framing through the casts of characters in the news. In *Proceedings of the 2016 conference on empirical methods in natural language processing*. 1410–1420.
- Paul DiMaggio, Manish Nag, and David Blei. 2013. Exploiting affinities between topic modeling and the sociological perspective on culture: Application to newspaper coverage of US government arts funding. *Poetics* 41, 6 (2013), 570–606.
- Caitlin Doogan and Wray Buntine. 2021. Topic model or topic twaddle? re-evaluating semantic interpretability measures. In *Proceedings of the 2021 conference of the North American chapter of the association for computational linguistics: human language technologies*. 3824–3848.
- James N Druckman. 2001. The implications of framing effects for citizen competence. *Political behavior* 23, 3 (2001), 225–256.
- Robert M Entman. 1993. Framing: Toward clarification of a fractured paradigm. *Journal of communication* 43, 4 (1993), 51–58.
- Elena Erosheva, Stephen Fienberg, and John Lafferty. 2004. Mixed-membership models of scientific publications. *Proceedings of the National Academy of Sciences* 101, suppl_1 (2004), 5220–5227.
- Norman Fairclough. 2013. *Critical discourse analysis: The critical study of language*. Routledge.
- William A Gamson. 1989. News as framing: Comments on Graber. *American behavioral scientist* 33, 2 (1989), 157–161.
- William A Gamson and Andre Modigliani. 1989. Media discourse and public opinion on nuclear power: A constructionist approach. *American journal of sociology* 95, 1 (1989), 1–37.
- Stuart Geman and Donald Geman. 1984. Stochastic relaxation, Gibbs distributions, and the Bayesian restoration of images. *IEEE Transactions on pattern analysis and machine intelligence* 6 (1984), 721–741.
- Erving Goffman. 1974. *Frame analysis: An essay on the organization of experience*. Harvard University Press.
- Thomas L Griffiths and Mark Steyvers. 2004. Finding scientific topics. *Proceedings of the National academy of Sciences* 101, suppl_1 (2004), 5228–5235.
- Maarten Grootendorst. 2022. BERTopic: Neural topic modeling with a class-based TF-IDF procedure. *arXiv preprint arXiv:2203.05794* (2022).
- Lars Guenther, Maria Gaertner, and Jessica Zeitz. 2021. Framing as a concept for health communication: A systematic review. *Health Communication* 36, 7 (2021), 891–899.

952	Kirk Hallahan. 1999. Seven models of framing: Implications for public relations. <i>Journal of public relations research</i> 11, 3 (1999), 205–242.	1007
953		1008
954		1009
955	Amin Hosseiny Marani. 2025. <i>What Are Topic Modeling Metrics Measuring? A Comparison of Topic Modeling Metrics and Human Assessment Approaches</i> . Doctoral dissertation. Lehigh University.	1010
956		1011
957		1012
958		1013
959	Amin Hosseiny Marani, Joshua Levine, and Eric PS Baumer. 2022. One Rating to Rule Them All? Evidence of Multidimensionality in Human Assessment of Topic Labeling Quality. In <i>Proceedings of the 31st ACM International Conference on Information & Knowledge Management</i> . 768–779.	1014
960		1015
961		1016
962		1017
963		
964		
965	Alexander Hoyle, Pranav Goel, Andrew Hian-Cheong, Denis Peskov, Jordan Boyd-Graber, and Philip Resnik. 2021. Is automated topic model evaluation broken? the incoherence of coherence. <i>Advances in neural information processing systems</i> 34 (2021), 2018–2033.	1018
966		1019
967		1020
968		1021
969		1022
970		
971	Carina Jacobi, Wouter Van Atteveldt, and Kasper Welbers. 2018. Quantitative analysis of large amounts of journalistic texts using topic modelling. In <i>Rethinking Research Methods in an Age of Digital Journalism</i> . Routledge, 89–106.	1023
972		1024
973		1025
974		1026
975		1027
976		1028
977	Justin Jager, Diane L Putnick, and Marc H Bornstein. 2017. II. More than just convenient: The scientific merits of homogeneous convenience samples. <i>Monographs of the society for research in child development</i> 82, 2 (2017), 13–30.	1029
978		1030
979		1031
980		1032
981	Daniel Kahneman and Amos Tversky. 1984. Choices, values, and frames. <i>American psychologist</i> 39, 4 (1984), 341.	1033
982		1034
983		1035
984		1036
985	Shima Khanehzar, Andrew Turpin, and Gosia Mikolajczak. 2019. Modeling political framing across policy issues and contexts. In <i>Proceedings of the The 17th Annual Workshop of the Australasian Language Technology Association</i> . 61–66.	1037
986		1038
987		1039
988		
989	Lingpeng Kong, Nathan Schneider, Swabha Swayamdipta, Archana Bhatia, Chris Dyer, and Noah A Smith. 2014. A dependency parser for tweets. In <i>Proceedings of the conference on empirical methods in natural language processing</i> . Association for Computational Linguistics, 1001–1012.	1040
990		1041
991		1042
992		1043
993		1044
994		1045
995	Jim A Kuypers. 2010. Framing analysis from a rhetorical perspective. <i>Doing news framing analysis: Empirical and theoretical perspectives</i> (2010), 286–311.	1046
996		1047
997		1048
998	Richard R Lau and Mark Schlesinger. 2005. Policy frames, metaphorical reasoning, and support for public policies. <i>Political Psychology</i> 26, 1 (2005), 77–114.	1049
999		1050
1000		1051
1001		1052
1002		1053
1003	Zongxia Li, Andrew Mao, Daniel Stephens, Pranav Goel, Emily Walpole, Alden Dima, Juan Fung, and Jordan Boyd-Graber. 2024. Improving the TENOR of Labeling: Re-evaluating Topic Models for Content Analysis. In <i>Proceedings of the 18th Conference of the European Chapter of the Association for Computational Linguistics (Volume 1: Long Papers)</i> . 840–859.	1054
1004		1055
1005		1056
1006		1057
		1058
	Yijia Liu, Yi Zhu, Wanxiang Che, Bing Qin, Nathan Schneider, and Noah A Smith. 2018. Parsing tweets into universal dependencies. <i>arXiv preprint arXiv:1804.08228</i> (2018).	
	John Lofland, David Snow, Leon Anderson, and Lyn H Lofland. 2022. <i>Analyzing social settings: A guide to qualitative observation and analysis</i> . Waveland Press.	
	Christopher Lucas, Richard A Nielsen, Margaret E Roberts, Brandon M Stewart, Alex Storer, and Dustin Tingley. 2015. Computer-assisted text analysis for comparative politics. <i>Political Analysis</i> 23, 2 (2015), 254–277.	
	Christopher D Manning, Mihai Surdeanu, John Bauer, Jenny Rose Finkel, Steven Bethard, and David McClosky. 2014. The Stanford CoreNLP natural language processing toolkit. In <i>Proceedings of 52nd annual meeting of the association for computational linguistics: system demonstrations</i> . 55–60.	
	Julia Mendelsohn, Ceren Budak, and David Jurgens. 2021. Modeling framing in immigration discourse on social media. <i>arXiv preprint arXiv:2104.06443</i> (2021).	
	Fred Morstatter, Liang Wu, Uraz Yavanoglu, Stephen R Corman, and Huan Liu. 2018. Identifying framing bias in online news. <i>ACM Transactions on Social Computing</i> 1, 2 (2018), 1–18.	
	Nona Naderi and Graeme Hirst. 2017. Classifying frames at the sentence level in news articles. <i>Pol- 9</i> (2017), 4–233.	
	David Newman, Jey Han Lau, Karl Grieser, and Timothy Baldwin. 2010. Automatic evaluation of topic coherence. In <i>Human language technologies: The 2010 annual conference of the North American chapter of the association for computational linguistics</i> . 100–108.	
	Zhongdang Pan and Gerald M Kosicki. 1993. Framing analysis: An approach to news discourse. <i>Political communication</i> 10, 1 (1993), 55–75.	
	Hyojung Park and Bryan H Reber. 2010. Using public relations to promote health: A framing analysis of public relations strategies among health associations. <i>Journal of Health Communication</i> 15, 1 (2010), 39–54.	
	Alexandrin Popescul, Lyle H Ungar, David M Pennock, and Steve Lawrence. 2013. Probabilistic models for unified collaborative and content-based recommendation in sparse-data environments. <i>arXiv preprint arXiv:1301.2303</i> (2013).	

1059	Forough Poursabzi-Sangdeh, Jordan Boyd-Graber, Leah	Laure Thompson and David Mimno. 2018. Author-	1112
1060	Findlater, and Kevin Seppi. 2016. Alto: Active learn-	less topic models: Biasing models away from known	1113
1061	ing with topic overviews for speeding label induction	structure. In <i>Proceedings of the 27th International</i>	1114
1062	and document labeling. In <i>Proceedings of the 54th</i>	<i>Conference on Computational Linguistics</i> . 3903–	1115
1063	<i>Annual Meeting of the Association for Computational</i>	3914.	1116
1064	<i>Linguistics (Volume 1: Long Papers)</i> . 1158–1169.		
1065	Marta Recasens, Cristian Danescu-Niculescu-Mizil, and	Amos Tversky and Daniel Kahneman. 1985. <i>The</i>	1117
1066	Dan Jurafsky. 2013. Linguistic models for analyz-	<i>framing of decisions and the psychology of choice</i> .	1118
1067	ing and detecting biased language. In <i>Proceedings</i>	Springer.	1119
1068	<i>of the 51st Annual Meeting of the Association for</i>		
1069	<i>Computational Linguistics (Volume 1: Long Papers)</i> .	Wouter Hendrick Van Atteveldt. 2008. Semantic net-	1120
1070	1650–1659.	work analysis: Techniques for extracting, represent-	1121
		ing, and querying media content. (2008).	1122
1071	Alan Ritter, Oren Etzioni, et al. 2010. A latent dirich-	Baldwin Van Gorp et al. 2010. Strategies to take sub-	1123
1072	let allocation method for selectional preferences. In	jectivity out of framing analysis. <i>Doing news fram-</i>	1124
1073	<i>Proceedings of the 48th Annual Meeting of the Asso-</i>	<i>ing analysis: Empirical and theoretical perspectives</i>	1125
1074	<i>ciation for Computational Linguistics</i> . 424–434.	(2010), 84–109.	1126
1075	Margaret E Roberts, Brandon M Stewart, Edoardo M	Dror Walter and Yotam Ophir. 2019. News frame anal-	1127
1076	Airolidi, K Benoit, D Blei, P Brandt, and A Spirling.	ysis: An inductive mixed-method computational ap-	1128
1077	2014. Structural topic models. Retrieved May 30	proach. <i>Communication Methods and Measures</i> 13,	1129
1078	(2014), 2014.	4 (2019), 248–266.	1130
1079	Michael Röder, Andreas Both, and Alexander Hinneb-	Tuukka Ylä-Anttila, Veikko Eranti, and Anna Kukkonen.	1131
1080	urg. 2015. Exploring the space of topic coherence	2022. Topic modeling for frame analysis: A study of	1132
1081	measures. In <i>Proceedings of the eighth ACM inter-</i>	media debates on climate change in India and USA.	1133
1082	<i>national conference on Web search and data mining</i> .	<i>Global Media and Communication</i> 18, 1 (2022), 91–	1134
1083	399–408.	112.	1135
1084	Peter Salovey and Duane T Wegener. 2003. Commu-	Caleb Ziems and Diyi Yang. 2021. To protect and to	1136
1085	nunicating about health: Message framing, persuasion	serve? analyzing entity-centric framing of police	1137
1086	and health behavior. <i>Social psychological founda-</i>	violence. <i>arXiv preprint arXiv:2109.05325</i> (2021).	1138
1087	<i>tions of health and illness</i> (2003), 54–81.		
1088	Dietram A Scheufele. 1999. Framing as a theory of me-	A Models Assessment: Summary Table	1139
1089	dia effects. <i>Journal of communication</i> 49, 1 (1999),	This table summarizes the key points being made	1140
1090	103–122.	by participants in assessing each of the models.	1141
1091	Dietram A Scheufele. 2000. Agenda-setting, priming,	B Interactive Interface for	1142
1092	and framing revisited: Another look at cognitive ef-	Human-Subject Model Evaluation	1143
1093	fects of political communication. <i>Mass communica-</i>		
1094	<i>tion & society</i> 3, 2-3 (2000), 297–316.	B.1 Interface for LDA	1144
1095	Donal Schön and Martin Rein. 1994. Frame reflection:	The LDA model includes three main components:	1145
1096	Toward the resolution of intractable policy controver-	topic terms, their probabilities, and example docu-	1146
1097	sies. <i>Basic Book</i> (1994).	ments in which the terms appear. Topic terms	1147
1098	Donal Shön and Martin Rein. 1994. Frame reflection:	probabilities, for LDA (in addition to the other two	1148
1099	Toward the resolution of intractable policy controver-	models), are visualized using a horizontal bar, in-	1149
1100	sies. <i>Basic Book</i> (1994).	stead of percentages. This design is informed by	1150
1101	Alison Smith, Tak Yeon Lee, Forough Poursabzi-	iterative feedback shared among authors, and with	1151
1102	Sangdeh, Jordan Boyd-Graber, Niklas Elmqvist, and	pilot testing with our lab-mates.	1152
1103	Leah Findlater. 2017. Evaluating visual represen-	Example documents are shown for each topic	1153
1104	tations for topic understanding and their effects on	term, provided that the example documents are	1154
1105	manually generated topic labels. <i>Transactions of the</i>	within the topic’s context (i.e., have a high prob-	1155
1106	<i>Association for Computational Linguistics</i> 5 (2017),	ability for the given topic). The interface shows	1156
1107	1–16.	the snippets of these example documents, in which	1157
1108	Mihai D Sturdza et al. 2018. Automated framing analy-	the associated topic term appears. This design was	1158
1109	sis: A rule based system for news media text. <i>Journal</i>	made ensure that the user was not overwhelmed	1159
1110	<i>of Media Research-Revista de Studii Media</i> 11, 32	with a copious amount of text at once. However,	1160
1111	(2018), 94–110.		

Criterion	LDA	LDA-GR	LLTR
Context	✓ Sufficient Contexts to capture discussed issues ✗ Falls short in providing the broader overview of corpus ✗ Lacks diverse contexts	✓ Sufficient contexts to capture discussed issues ✗ Falls short in providing the broader overview of ideas discussed in corpus ✗ Lacks diverse contexts	✓ Provided the most effective contexts by capturing co-occurring terms ✓ Context readily and easily available through offering co-occurring terms ✓ Offered contexts were diverse and comprehensive
Clarity	✗ Lacks clear connections between example documents ✗ Required a lot of reading of the documents in full to clarify the meaning of words ✗ Required a lot of reading of the documents in full to confirm the inferred framing	✓ Grammatical relationships helped clarify how the topic terms are used ✗ Efforts required to account for grammatical relationships made this process less effective ✗ Required a lot of reading to confirm inferred framing	✓ Made meaning of topic terms clear, due to providing their co-occurring terms ✓ Easy clarification process, due to providing the broader overview of the corpus
Confidence	✗ Lack of diverse contexts, and sparsity of context supporting each framing evidence made researchers less confident about the results being representative ✗ Less confidence about representatives of results reduced participants' confidence in their analysis	✗ Offered more confidence about connecting the topic terms, due to providing grammatical relationships. ✗ Offered less confidence about whether model results are representative of the broader corpus	✓ Offered confidence in model's results being representative of the broader overview of the corpus. ✓ Made participants more confident about their framing analysis
Curve	✓ Easiest model learning curve, due to reduced numbers of components ✗ Made it difficult to infer framing, due to the lack of supportive components to find connection between components	✗ Increased model learning curve, due to the addition of grammatical relationships ✗ Difficult to find connections between example documents	✗ The steepest learning curve, due to increased components ✓ Once passed learning curve, it was easier to find framing evidence

Table 1: Comparison of the LDA, LDA-GR, and LLTR models in terms of context, clarity, confidence, and curve. The LLTR model provided the most diverse and interconnected contexts, enhancing the clarity of framing evidence and resulting in the highest confidence in model results, thereby participant's highest confidence in their own framing analysis. However, LLTR requires the steepest learning curve.

participants are instructed that they can view the full example documents upon clicking on these snippets.

To help emphasize how each topic term appears in the example document, the topic term is bolded in its associated example documents. This visual cue aims to facilitate comprehensions.

Figure 2 shows a screenshot of the LDA model interface. Topic terms are sorted from most to least probable, with their probabilities visualized as horizontal bars.

B.2 Interface for LDA-GR

The LDA-GR model’s interface is built to be very similar to the LDA model, with addition of the grammatical relationship component. Specifically, the LDA-GR interface includes four components, topic terms, their probability, the grammatical relationship in which each topic term appears, as well as example documents in which each pair of topic term and its captured grammatical relationship appear, within the context of topic.

The grammatical relationship associated with each topic term is depicted using a subscript underneath the topic term (See Figure 4). Recognizing that these grammar terms are not part of everyday language, we added their definition in the interface. However, to avoid overwhelming the participants at the first glance, this functionality would only appear upon click. Specifically, readers can click on these grammatical relationships to see their definitions (e.g., “nominal subject”: noun or noun phrase that performs the action of the verb. Example: *Clinton* defeated Dole, wherein “Clinton” is the nominal subject because it is the entity performing the action of defeating in the sentence.), and how they appear in documents (See figure 3). Participants were instructed about these steps in the instruction section. See appendix C).

Piloting this model with a number of researchers, we learned that providing these grammatical relationships without their context makes it hard to for participants to make sense of. Therefore, in a post-processing step, we captured example documents in which these topic terms appear in their associated grammatical relationships, provided that those documents are representative of the topic (i.e., provided that the example document have a high probability for the given topic). Additionally, similar to the LDA interface, here, participants can view these example documents in full upon clicking on

these shown snippets.

Figure 4 depicts a screenshot of the interface of the LDA-GR model.

B.3 Interface for LLTR

The LLTR model includes four main components, topic terms, topic terms probability, their co-occurring terms, and example document in which topic term co-occurring with other topic terms within topic document. Similar to the two other models, the selection of these example documents is conditioned on having topic probability above a certain threshold for the given topic(i.e., threshold = 0.40). While grammatical relationships are considered when capturing linked theta roles and influence co-occurring terms for topic terms, they are not explicitly displayed in the experiments to avoid an overly complex interface. This design is formed by the iterative feedback received in piloting the interface. Additionally, similar to the interface for the other two models, participants can view the example documents upon clicking on these shown snippets.

For example, for the topic term *support* (See Figure 5), the relative probability of this term is visualized using a horizontal bar (this bar decrease for terms with lower probability for the given topic), a set of its co-occurring words, including *provide*, *need*, *services*, *families*, etc. For each pairs of topic term *support* and ech of its co-occurring words (e.g., *provide*), the example column shows multiple examples in which these terms co-occur together in the topic context. Due to the inclusion of co-occurring terms in the LLTR model, showing all these example documents at the landing page would make the interface even more complex. To address this concern, the multiple example documents are collapsed, and participants were instructed to click the expansion button, shown with a “+” sign, to see more examples documents for each pairs of topic term and its co-occurring words.

Figure 5 depicts a screenshot of the interface of the LLTR model.

C Study Instructions for Human-subject Study

This section provide an overview of the consent, and the instructions offered to the participants in the human-subject study.

The initial step involved participants reviewing the consent form. This document outlined the

Probability	Term	Examples
	support	I understand there will be some who need to travel from other states to return to a home in Vermont or support a vulnerable family member. ... These are challenging times, and we must support one another, not take advantage of others, said Governor Whitmer. ... These are challenging times, and we must support one another, not take advantage of others, said Governor Whitmer. ...
	businesses	In the interest of public health, we are requiring modifications in operations for businesses that serve food and drinks, and temporarily prohibiting interstate games and tournaments for indoor K-12 sports. ... 230, which will increase indoor capacity limits for certain businesses and increase both the general indoor and outdoor gathering limit. ... Tony Evers today announced another turn of the dial on Safer at Home to add even more opportunities for Wisconsin businesses to get back to work in a safe and responsible way. ...

Figure 2: A screenshot of the LDA model's interface, which includes topic terms, their probability, and the example document in which they appear. Note: this screenshot only presents part of the topic, to give an overview of the the model components, while ensuring concision.

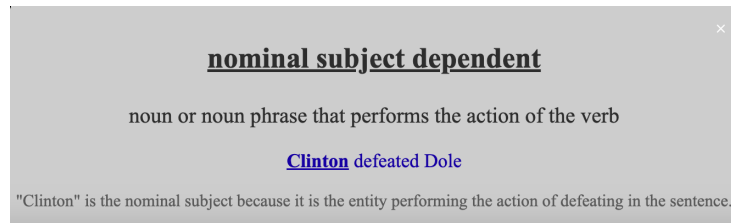


Figure 3: The screenshot of an example of grammatical relationship definitions displayed upon clicking within the LDA-GR interface.

Probability	Term	Examples
	governor _(noun compound modifier - dependent)	sacramento - governor gavin newsom and state health officials will hold a media availability today to provide an update on the states response to covid-19. ... sacramento - governor gavin newsom will provide an update tomorrow on the states response to wildfires and the covid-19 pandemic. ... sacramento - governor gavin newsom will provide an update tomorrow on the states response to the covid-19 pandemic. ...
	order _(noun compound modifier - governor)	licensees multiple violations of the current michigan department of health and human services (mdhhs) emergency order include : allowing non-residential , in-person gatherings ; providing in-person dining ; failure to require face coverings for staff and patrons ; and failure to prohibit patrons from congregating. ... executive order 2020-109 , which takes effect immediately and continues through june 12 , 2020 , extends the following health and safety guidelines , among others : executive order 2020-108 which also takes effect immediately and continues through june 26 , 2020 — maintains restrictions on visitation to g health care facilities , residential care facilities , congregate care facilities , and juvenile justice facilities , but authorizes the department of health and human services to gradually re-open visitation as circumstances permit denver , june 4 , 2020 : in accordance with governor jared polis executive order and public health order 20-28 , , the colorado department of public health and environment today finalized guidance outlining the steps required to allow personal and outdoor recreation activities to resume while minimizing the potential spread of covid-19

Figure 4: A screenshot of the LDA-GR model's interface, which includes topic terms, their probability scores, the grammatical relationship in which they appear, and the example documents of the appearance of each topic term in its associated grammatical relationship within the corpus. Note that this screenshot only presents part of the results, to give an overview of the the model' components, while ensuring concision.

Probability	Term	Co-occurring	Example
	support	provide	+ , beginning january 23 and throughout the severe weather, thetexas division of emergency management to provide support to localjurisdictions and conduct preliminary damage assessments in coordination withlocal officials ...
		need	+ businesses in our state have experienced immense challenges since the covid-19 pandemic began, and they need our support, governor kelly said ...
		services	+ employment and training services (\$7 million grant) - this funding will expand career support services supported by the workforce investment boards throughout the state ...
		families	were going to continue working to make sure that every wisconsinite knows how these funds are being used to fight the pandemic and support families , farmers, and small businesses who need it most ...
		programs	weigand will help guide the states pandemic response and support agency programs in a post-pandemic ohio to develop modern, innovative approaches to address all public health needs ...
		businesses	this bill will give our restaurants more certainty for the future so they can once again lean into the outdoor expansions we allowed this past summer to help recoup losses and strengthen their businesses and the jobs they support ...
	today	signed	+ boise, idaho - governor brad little signed an executive order today , forming his new coronavirus financial advisory committee to oversee the approximately \$1 ...
		released	+ governor jay inslee released a statement today regarding the announcement of president joe bidens american jobs plan, the first part of his build back better agenda ...
		issued	+ sacramento - governor gavin newsom issued the below statement today following the houses passage of the american rescue plan:i applaud president biden and speaker pelosi on the passage of the american rescue plan - \$1 ...

Figure 5: Screenshot of the LLTR model interface, including topic terms, their probability, a set of co-occurring terms for each topic term, and example documents in which each topic term appears with its co-occurring terms. Note that this screenshot only presents part of the results, to give an overview of the the model' components, while ensuring concision.

study's objectives, procedures, data handling, and compensation, which consisted of a \$50 Amazon gift card. Participants were then given the choice to withdraw or to continue by indicating their consent. Only if they consented to participants, they were able to proceed with the rest of the survey.

Before assigning participants to one of the randomized experimental conditions, as described in the paper, participants were provided with the definition of framing, as described below.

What is Framing? Framing is a dynamic and constantly evolving set of processes by which people construct their understanding of the world's events. The processes of framing help organize facts and information to give them meaning. Framing influences our understanding both of major world events, such as the COVID-19 pandemic, and of our personal daily experiences, such as a visit to the doctor's office.

Framing involves different processes, including the following:

- Determining what counts as an issue.
- Diagnosing the causes of those issues.
- Making moral judgments, such as about what is right and wrong, or about how people ought to behave.

- Suggesting potential remedies to address the issues.

One way of understanding framing is by closely examining the patterns of language that people use to describe a situation. This language does not occur in only one piece of content at a time, such as a single news story or a single social media post. Rather, framing happens through consistent patterns of word use across different pieces of content.

Participants were asked to spend at least *one minute* reading the above definition, before they could see the next button to proceed.

Next, further instructions was provided to explain the study's main tasks, described below.

Instructions on the main tasks: Please read the following instruction very carefully as it provides a better sense of the study's goal, and how to complete the study's steps:

In the following pages, you will be asked to review the results from two models, named Alpha and Beta, respectively, and explore framing processes using these results.

As mentioned earlier, framing occurs across different pieces of content. However, reading over all the pieces of content to examine framing can be really

time consuming, if not impossible. To help address this challenge, these two models were designed to help extract patterns of language across a very large set of documents to help researchers in exploring and understanding framing process in a corpus of almost 4,000 news releases from the Department of Health in different states during the COVID-19 pandemic.

Both of these models identify topics or groups of thematically related words that co-occur. However, each model identifies topics in a slightly different way.

Your task is to leverage these results and evaluate these models in terms of their utility in exploring framing process across large corpus of documents. Specifically, you will investigate framing processes based on the results provided using each of these models, the alpha model and then the beta model, and later evaluate the utility of each of these models.

Once participants read over the main instructions, they were prompted about the questions that they were going to respond to after reviewing the results. This instruction is described in the following passages.

Prompts about the survey questions: In this section, you will review three topics from our first model, named Alpha.

After exploring all the three topics, you will be asked to respond to some questions about framing processes, as evidenced within and across these topics.

These questions are listed below:

What are the issues that are being discussed?

What are the suggested or implied causes of those issues? What moral judgments are being made related with those issues? What potential remedies are being suggested? Are there any other observations you make about these topics and its relation with framing?

Then, after working with each of the models, you will be asked to evaluate the

model's utility in helping you respond to these above questions.

Note: Please feel free to take notes while you are examining these results to refer to when are going to attend to the above questions. You can also leave the links open to refer back to and review as needed.

At this step participants were randomly assigned to one of experimental conditions. The instructions for each of the model were provided as follows.

C.1 The LDA Instruction

In the following section, you will review three topics using this model explore framing processes.

This model that accounts for patterns of word co-occurrence. We would like you to go over three topics captured using this model, and investigate framing evidence using each of these topics.

The picture below shows a screenshot of the model's interface showing each component of the model's interface: First, there is a list of (1) the top terms for the topic. For each top term, it also shows (2) the probability of the term for that topics. The interface also includes (3) top documents, by clicking on each document, (4) the document's full text will appear in a box.

At this step, we encouraged participants to take notes while they were examining these results to attend to the questions provided to them earlier.

C.2 LDA-GR Instruction

In the following section, you will review three topics captured using this model and explore framing processes evidence in these topics.

This model that accounts for patterns of word co-occurrence, as well as the grammatical relationships in which words occur.

Picture below shows a screenshot of the model's interface, and call out different components of the results: First, there is a list of (1) the top terms for the topic. For each top term, it also shows (2) the probability of the term for that topics. The interface also shows (3) the grammatical relationship in

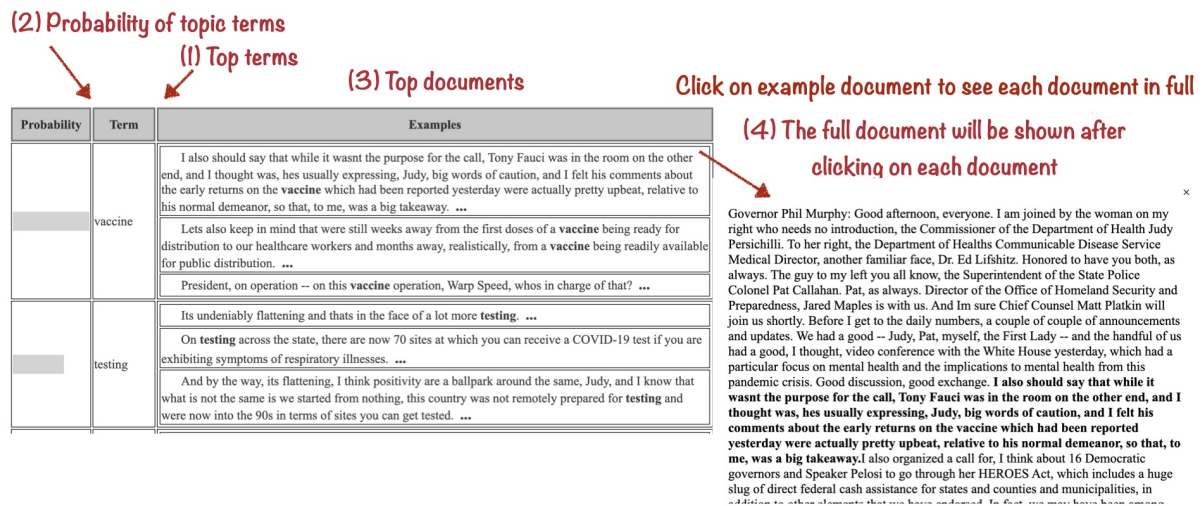


Figure 6: The LDA model's demo, shown to the participants before they started using the model's interface.

which topic terms occur. By clicking on the grammatical relationship associated with each topic term, you will see a more detailed (4) description of the grammatical relationship. In front of each term, you see a list of (5) example document in which the topic term appears. By clicking on each of these documents, (6) the document's full text will appear in a box.

C.3 LLTR Instruction

In the following section, you will review three topics using this mode to explore framing processes.

This model captures two components. The first part captures groups of thematically related words that tend to occur in documents together. For each of those words, the model also identifies the other words directly related with it, e.g., as the subject of a verb, or an adjective modifying a noun.

However, instead of capturing a one to one link between words and their grammatical relationship, it captures the distribution of grammatical relationships in which terms occur as well as the other terms by which each term co-occur in a relationship. To simplify the results, we extracted and demonstrated example documents in which pairs of words co-occur in a grammatical relationship.

The picture below shows a screenshot

of the model's interface, showing each component of the model and the way results are organized. First, there is a list of (1) the top terms for the topic. For each top term, it also shows (2) the probability of topic term, (3) the occurring terms that appear with each topic terms, as well as the (4) example documents in which these terms co-occur. By clicking on the example documents, you can see (5) the full text of the example document. Some pairs of topic terms and co-occurring words have more than one example. Clicking on (6) the + sign to see the list of more examples, and click on each to see the examples' full text.

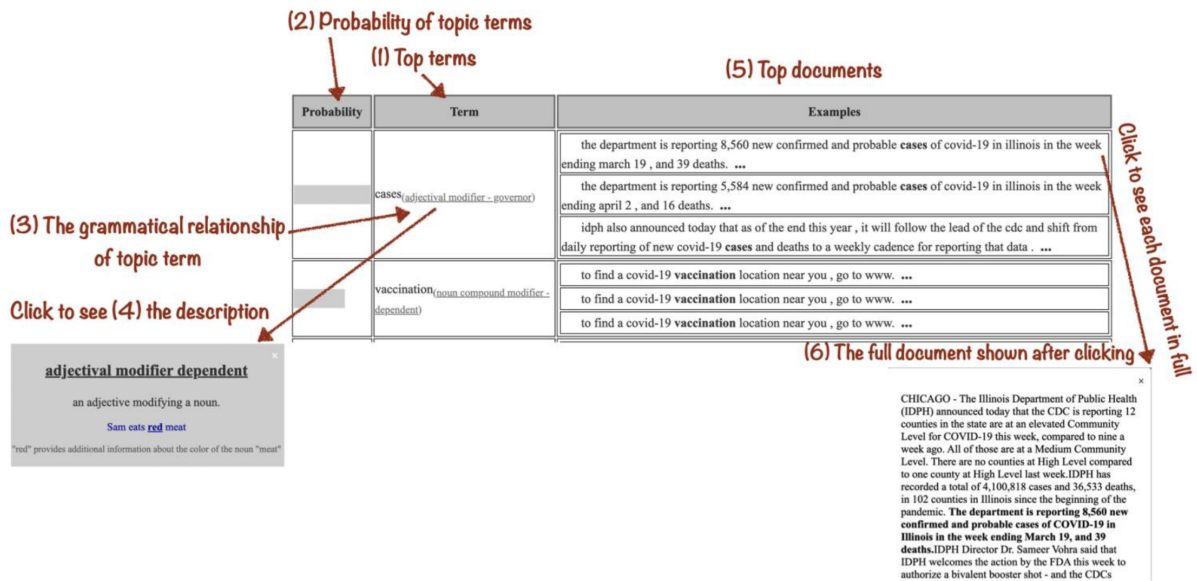


Figure 7: The LDA-GR model's demo, shown to the participants before they started using the model's interface.



Figure 8: The LLTR model's demo, shown to the participants before they started using the model's interface.