

Article

Graph-Based Lexical Sentiment Propagation Algorithm

Tajana Ban Kirigin ^{1,*}, Sanda Bujačić Babić ^{1,*} and Benedikt Perak ^{2,†}¹ Faculty of Mathematics, University of Rijeka, 51000 Rijeka, Croatia² Faculty of Humanities and Social Sciences, University of Rijeka, 51000 Rijeka, Croatia; bperak@uniri.hr

* Correspondence: bank@uniri.hr (T.B.K.); sbujacic@uniri.hr (S.B.B.)

† These authors contributed equally to this work.

Abstract: In the rapidly developing field of sentiment analysis, it is a challenge to create sentiment dictionaries with broad coverage, especially for languages with limited resources. To address this problem, we propose innovative methodologies that automate the creation of comprehensive sentiment dictionaries, utilising both traditional linguistic approaches and state-of-the-art artificial intelligence technologies. The methodologies are characterised by their universal applicability to different languages. The proposed ConGraCNet Sentiment Propagation algorithm uniquely combines existing sentiment dictionaries and corpus-based syntactic–semantic embedding graphs to reliably capture and propagate sentiment values in lexical networks. To demonstrate the particular benefit for underrepresented languages with scarce sentiment resources, such as Croatian, we used the ConGraCNet Sentiment Propagation algorithm to create the *Sentiment-hr* dictionary and the AI tool GPT to generate the *Sentiment-hr-AI* dictionary. The two open-source sentiment dictionaries created are the largest and most comprehensive resources for the Croatian language to date, being at least ten times larger than the second-largest sentiment dictionary available. Our results demonstrate the effectiveness of the methods presented, which significantly expand the toolkit of sentiment analysis for the Croatian language and provide researchers with valuable insights and resources.

Keywords: sentiment analysis; sentiment dictionaries; applications of graph data processing; complex networks; algorithmic sentiment propagation; AI-driven sentiment propagation

MSC: 05C90; 68R10

Received: 28 February 2025

Revised: 27 March 2025

Accepted: 29 March 2025

Published: 31 March 2025

Citation: Ban Kirigin, T.; Bujačić Babić, S.; Perak, B. Graph-Based Lexical Sentiment Propagation Algorithm. *Mathematics* **2025**, *13*, 1141. <https://doi.org/10.3390/math13071141>

Copyright: © 2025 by the authors. Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<https://creativecommons.org/licenses/by/4.0/>).

1. Introduction

One of the most remarkable achievements in the field of computational linguistics and natural language processing (NLP) is the ability of algorithms to analyse and interpret textual data comprehensively. The branch of NLP known as sentiment analysis focuses on determining the affective tone or sentiment of the input, using computer-assisted methods to analyse the affective components in spoken and written language [1,2].

Sentiment analysis involves the systematic identification and extraction of affective content such as polarity (positive, negative or neutral) and emotions from texts. This also includes the classification and/or assignment of a normalised range of values to dimensions such as hedonic valence [3]. Sentiment can be analysed for words, concepts, multi-word phrases, sentences, paragraphs or entire texts.

In the growing research field of analysing emotions in texts, sentiment dictionaries are an important resource for the development of automatic sentiment analysis systems [4–10]. Sentiment dictionaries are collections of words categorised as either positive, neutral or

negative, or assigned numerical values that estimate the intensity of affective charge within a particular class or on a selected scale, typically a scale from -1 to 1 , with -1 being the most negative and 1 being the most positive affective value. The sentiment values of individual lexemes are used as the basis for more complex evaluations of larger linguistic forms such as phrases, sentences or entire texts.

A significant challenge in employing sentiment dictionaries lies in the subjective nature of determining the emotional significance of a given word [11]. This difficulty is further compounded by linguistic phenomena such as polysemy and homonymy, where a single word can have multiple meanings, sometimes unrelated, yet these meanings are often categorised under a single dictionary entry [12–14]. Also, some words may have multiple part-of-speech or other grammatical tags, e.g., *work-n* vs. *work-v*, *fight-n* vs. *fight-v* which is not always classified in the dictionary.

Another major problem with sentiment dictionaries is their low coverage, especially in languages with fewer resources [15]. In linguistic contexts where sentiment dictionaries are under-resourced, lexicons are often translated into English, and then the sentiment scores from an English-language sentiment dictionary are applied. Similarly, multilingual sentiment lexicons [16] can be referenced and used in AI approaches to achieve multilingual generalisation in language model pretraining [17]. However, this approach often leads to suboptimal results, as the sentiment scores assigned in this way may not accurately capture the semantic and affective nuances of the original terms in their native language. In English, for example, the lexeme *adventure-n* usually has a positive or exciting connotation. In Croatian, however, the corresponding lexeme *avantura-n* can also refer to a reckless or irresponsible act, particularly in the context of personal relationships (similar to eng. *affair*). In direct translation, sentiment analysis models based on English lexicons cannot capture the potentially negative connotation of such lexemes, resulting in a very different value of the sentiment or its misclassification.

This paper presents an algorithmic approach that aims to improve linguistic resources by computationally augmenting existing sentiment dictionary entries. Using the sentiment values available in sparse dictionaries, the method projects the emotional scores onto an extensive collection of semantically related words. The result is a comprehensive sentiment dictionary with a significantly larger coverage, expanding the scope of sentiment analysis in the target language.

This proposed algorithmic method integrates interdisciplinary concepts, drawing theoretical foundations from computational linguistics, cognitive linguistics and corpus-based linguistics. It extends the principles of Construction Grammar Conceptual Networks (ConGraCNet)—an established graph-based method developed for the extraction of semantically associated lexical networks [18,19]. The ConGraCNet approach utilises a linguistically annotated corpus and exploits syntactic and semantic dependency annotations that are crucial for improving the understanding of natural language and sentiment analysis algorithms. These include accurately resolving the polysemy of lexemes [12], effectively annotating word senses [20], and calculating the sentiment potential inherent in polysemous terms [21]. This methodology has been effectively applied in the ConGraCNet framework of the EmocNet project [22] and represents a fundamental element of our theoretical approach to the challenges and key contributions addressed in this work.

In addition to the aforementioned algorithmic approach, we also introduced a novel methodology for constructing a sentiment dictionary using large language models (LLMs).

Our approach facilitates comprehensive sentiment analysis and serves as a valuable tool for comparing the results of different sentiment analysis techniques. By integrating this AI-generated sentiment dictionary, we expect to significantly refine and improve sentiment classification techniques. This integration offers the opportunity to explore the synergy

between human-generated sentiment analysis and state-of-the-art artificial intelligence and promises a more robust and versatile toolkit for researchers and practitioners in the field.

The main contributions of this work are outlined below:

1. **ConGraCNet Sentiment Propagation Algorithm:** The algorithm for the automatic generation of a broad-coverage sentiment dictionary for a selected language based on an existing sentiment dictionary and a corpus-based syntactic-semantic embedding graph. This contribution is particularly important for the study of sentiment analysis of languages for which available sentiment dictionaries have low coverage. It applies to most languages due to the universal representation of semantic networks; It is a transparent and easily explainable traditional approach.
2. **Syntactic–Semantic-hrWac Embedding Graph [23]:** A lexical graph structure constructed utilising the hrWac corpus [24] with the application of the ConGraCNet methodology. This graph is instrumental in mapping the graph structure of lexeme-centric networks, which are pivotal for the sentiment propagation algorithm. It facilitates the systematic propagation of sentiment values across such networks and provides a structured framework for the analytical examination of semantic domains within the corpus;
3. **Sentiment-hr dictionary [25]:** a sentiment dictionary for the Croatian language propagated using the hrWac Coordination Graph and sparse Croatian sentiment dictionary from BabelSenticNet [26]. It is currently the most comprehensive sentiment dictionary for the Croatian language and available as an open-access resource;
4. **Sentiment-hr-AI sentiment dictionary [27]:** This dictionary for the Croatian language has been constructed using OpenAI’s GPT-4 [28]. The creation of Sentiment-hr-AI represents a methodological advancement in the field, as it utilises the extensive natural language understanding capabilities of state-of-the-art LLMs. The primary aim of this dictionary is to facilitate comparative and methodological analysis within sentiment analysis research.

This paper is organised as follows. The next section provides a brief overview of the related approaches and available resources for sentiment analysis. In Section 3, we describe graph-based methods and algorithms for lexical sentiment propagation. Section 4 presents the propagation of a comprehensive Croatian sentiment dictionary based on the above-mentioned algorithm. In Section 6, we discuss the presented approach and compare the results of a propagation-driven approach using a sparse BabelNet Sentic dictionary with a sentiment dictionary built by the iterative process of extracting sentiment values using the large language model GPT-4. We conclude in Section 7, where we also make suggestions for future work.

2. Related Approaches and Available Resources

Linguistically, sentiment analysis involves the systematic study of affective states and subjective information [29–31]. This complex interdisciplinary field overlaps with computational linguistics, text analysis, and data mining, drawing its foundations from early work in computational linguistics and, more recently, from the rapid advancements in artificial intelligence [1,2]. Within this context, sentiment analysis is approached from two primary perspectives: lexicon-based methods and machine learning [32,33]. Hybrid approaches that integrate the latest deep learning strategies with lexicon-based and machine learning methods are increasingly being introduced [34].

Sentiment dictionaries play a central role in sentiment analysis, as they provide pre-defined sentiment values to words and phrases and thus serve as an important basis for algorithms to recognise and categorise sentiment in text data. These resources are significant

not only for their direct application in identifying the sentiment of specific lexemes, which is particularly relevant in rule-based or hybrid methods of sentiment analysis, but also as a means of facilitating the automatic assessment of the sentiment polarity of text. This capability is particularly beneficial for applications that are limited by the computational resources required for more sophisticated deep learning models [2,35,36].

Furthermore, sentiment dictionaries extend machine learning models by either contributing features in supervised learning environments or acting as essential components in unsupervised learning frameworks. This integration contributes to a deeper contextual understanding and enriches the models' ability to interpret nuances in sentiment [37,38].

Despite the inherent value of sentiment dictionaries in the realm of sentiment analysis, the process of creating, maintaining and developing these resources presents numerous challenges. These include the subjective nature of sentiment itself, linguistic complexities such as polysemy (words with multiple meanings) and homonymy (words with the same spelling or pronunciation but different meanings), and the dynamic nature of language evolution. Historically, these issues have contributed to the fact that there are few comprehensive sentiment and emotion lexicons.

These lexicons often reach their limits both in terms of the breadth of the lexical units they cover and the depth of the emotional dimensions they capture [39,40]. Overcoming these obstacles is crucial for the further development of sentiment analysis technologies and the expansion of their application to a variety of areas.

Within the spectrum of existing sentiment dictionaries, there is a notable diversity in the quantification and categorisation of feelings. Some dictionaries simply categorise words as positive or negative, offering a binary perspective on sentiment. In contrast, other dictionaries assign numerical values to lexemes, providing a more nuanced representation of sentiment intensity on a scale where positive numbers represent positive sentiments, negative numbers represent negative sentiments, and zero represents the absence of sentiment [41]. The numerical approach allows for a more detailed analysis of the sentiment intensity and therefore a more precise interpretation of the emotional tone in the text.

The traditional approach to creating sentiment dictionaries typically involves human annotators meticulously assigning sentiment labels, categories or numerical values to words and phrases based on their interpretation of the emotional tone conveyed by the language. This manual process, exemplified by the development of SentiWordNet [38]—a human-annotated lexicon of more than 115,000 entries derived from the extensive WordNet database—highlights the indispensable role of human expertise in distinguishing and categorising linguistic sentiment. Such methodologies, while labour-intensive, have laid the foundation for reliable sentiment analysis by leveraging the insights of human linguistics in conjunction with systematic methodologies for sentiment evaluation [42,43].

Advances in computational linguistics have led to more sophisticated methodologies for the expansion of sentiment dictionaries, involving a mixture of manual and automated processes. This modern approach usually involves two stages: First, a collection of seed sentiments is established either by manual annotation or by drawing from pre-existing dictionaries. Then, these seed values are algorithmically propagated over a basic graph of words, phrases or conceptual structures, expanding the sentiment dictionary in a systematic and scalable way [10,44–47]. This dual-step process is an example of the integration of human judgement and computational efficiency that facilitates the growth of sentiment resources.

Among the advanced resources in the field of sentiment analysis, SenticNet [10] is a comprehensive sentiment knowledge base customised for English and other languages, covering over 300,000 lexical concepts. This achievement is due to the seamless fusion of top-down and bottom-up learning methodologies, utilising a range of symbolic and

subsymbolic AI tools. At the heart of SenticNet is the Hourglass of Emotions model, which innovatively applies biologically inspired and psychologically motivated principles based on Plutchik's foundational work on human emotions. This model categorises emotions in different dimensions: polarity value, polarity intensity, introspection, temper, attitude, and sensitivity—enabling a multifaceted analysis of sentiment that skilfully combines computational intelligence with profound insights into human emotional states [48–50].

The integration of such advanced NLP and AI-driven methodologies represents a significant departure from the era of manual annotation in sentiment analysis. This evolution is particularly important when it comes to addressing the unique challenges of low-resourced languages that have historically been underserved by sentiment analysis research.

The launch of CroSentiLex [51] in 2012 was a pivotal moment for Croatian sentiment analysis, as it provides a corpus-based lexicon with an extensive collection of positive and negative words. This development and the publication of the first lexical dataset encoding emotions for the Croatian language in 2019 provided researchers and practitioners with crucial tools for sentiment analysis [52,53]. However, the reliance on the translation of Croatian texts into English for sentiment analysis reveals a critical quality gap in the direct consideration of the linguistic specificities of Croatian and thus emphasises the need for advanced, language-specific sentiment analysis tools [19,51].

Research in semi-supervised lexicon development and the application of computer modelling have further enriched the landscape of sentiment analysis. Early efforts focused on automating sentiment detection in Croatian financial texts [54], while later studies introduced semi-supervised methods for lexicon development using techniques such as latent semantic analysis and graph-based propagation [51]. These approaches highlight the potential of computational methods as a complement to manual expertise, particularly for languages with limited specialised resources.

In addition, advances in deep learning and the adaptation of models for cross-lingual sentiment analysis have shown that it is possible to apply sophisticated computational techniques to Croatian sentiment analysis. The use of convolutional neural networks and the development of datasets with sentiment labels for Croatian social media content, especially in response to the COVID-19 pandemic, emphasise the dynamic nature of sentiment analysis research. Comparative studies with multi-task models and the application of “zero-shot” and “few-shot” learning illustrate the versatility and effectiveness of advanced computational models in processing and understanding sentiment in Croatian texts [55–57].

Research comparing the effectiveness of word embeddings and string kernels in classifying sentiment has shown that word embeddings are superior in capturing the subtleties of sentiment, particularly in informal texts. This result underlines the importance of using modern computational techniques to improve sentiment analysis in Croatian, which is a language characterised by unique linguistic nuances [58]. There are sentiment lexicons for related Slavic languages, including the Czech, Macedonian, Polish, Slovakian, Slovenian, and Bosnian. These collective advances emphasise the significant progress made in the development of sentiment analysis tools for Croatian and other low-resourced languages but also highlight the continued need for innovative solutions.

3. Augmenting Sentiment Lexicons: Leveraging Graph Theory for Enhanced Dictionary Coverage

To mitigate the challenge of insufficient coverage in sentiment dictionaries, our methodology introduces an algorithm designed to enhance existing dictionaries by propagating sentiment values through an expanded lexicon. Central to our approach is the development of a lexical network that can be constructed from corpus-based syntactic dependencies [12,19] or other types of comprehensive network representations of synonymous lexical

relations. By using this network as a semantic embedding, we apply a sentiment value propagation mechanism that extracts the semantic similarity of lexemes and then transfers the sentiment value from neighbouring nodes to the root node. This process not only expands the scope of the dictionary but also enriches its multidimensionality with lexical graph structures that provide insight into the polysemous nature of lexemes. This strategy utilises the interconnected structure of lexical synonym relations, which are represented as graph objects, to systematically distribute sentiment values and thus close gaps in the coverage of the sentiment dictionary.

This lexical network can be built from a variety of sources in order to capture a comprehensive range of lexical relations. In particular, it can be built directly from a corpus by using coordinated syntactic dependencies (e.g., and/or constructions) to extract synonymous lexical relations. Furthermore, such relations can be derived from structured lexical databases such as WordNet, where explicit synonymous phrases provide a rich source of semantic associations. In addition, the advent of large language models provides a new way to extract synonym relations by utilising the extensive training of models on large text corpora to identify words and phrases with similar meanings. This multi-layered approach to building lexical network construction ensures a depth and breadth of semantic relationships that provide a robust basis for the propagation of sentiment scores, improving the coverage of the lexicon and its utility for sentiment analysis.

In the remainder of this section, we present the most important steps in the construction of the syntactic–semantic embedding lexical network for a given lexeme and the subsequent propagation of sentiment values.

3.1. Coordination-Based Syntactic-Semantic Embedding Lexical Graph

The propagation of sentiment values from the coordination of syntactic–semantic embedding lexical graphs is anchored in the ConGraCNet methodology [19]. The computational implementation of ConGraCNet is accessible via the GitHub repository [18] and involves several tasks, including the creation of tagged corpora, data retrieval from digital corpora, modelling, storage, algorithmic processing, sentiment analysis and the visualisation of syntactic–semantic structures. Based on the theory of construction grammar [59–62], this method assumes that coordination constructions [LEXEME_A and/or LEXEME_B] imply conceptual relatedness, which facilitates analyses of conceptual similarity, lexical ambiguity, semantic domain relatedness and sentiment. Due to the almost universal use of coordinated constructions and logical connectives in natural languages, the method is appropriate for the study of most languages.

To build a lexical graph from large corpora, a systematic approach focussing on statistical measures is used to identify the most relevant collocates for a given lexeme. This process involves several important steps that ensure that the created graph accurately reflects the semantic and syntactic relationships [19]. The method enables the calculation of multiple sentiment values for individual lexemes, thus facilitating the expansion of the lexicon to different sentiment dictionaries and their respective categories. Consequently, this approach leads to a significantly expanded sentiment dictionary that is characterised by its comprehensive coverage and multidimensional analysis capabilities.

3.2. Analysing Semantic Contexts: The Role of Lexical Networks in Lexical Graph Embeddings

Lexical networks play a crucial role in our sentiment propagation methodology, as they allow for targeted analysis around a central lexeme *a*. These networks, which are effectively subgraphs of the embedding lexical graph, focus on the immediate and extended network of relations of a single lexeme, emphasising its connections within the semantic landscape. The advantage of such networks lies in its ability to trace how sentiment values are influ-

enced by the semantic context of a lexeme. By focussing on these individual networks, our approach gains precision in the propagation of sentiment and ensures that the sentiment values assigned to lexemes reflect their specific semantic associations and usage contexts. This targeted analysis not only improves the accuracy of sentiment mapping but also deepens our insight into the intricate semantic networks that structure language.

The construction of a lexical network for a particular lexeme a involves the selection of vertices that belong to the same part-of-speech class (nouns, adjectives, adverbs or verbs) and have a high logDice value. Using the syntactic and semantic properties inherent in the coordination relationship $[LEXEME_A \text{ and/or } LEXEME_B]$, we generate a second-order friend-of-a-friend FoF_a network for a particular lexeme a . This resulting lexical network, FoF_a, contains lexemes that are semantically associated with a , thus facilitating the identification of lexemes that share prototypical conceptualisations within the same semantic domain.

To enable researchers and practitioners to construct and analyse such networks for their own purposes, a Python script is provided that automates this process. This script and documentation can be downloaded from the GitHub repository [63] which contains the necessary tools to run the embedding graph and generate customised lexical FoF_a networks for any lexeme a . This resource is designed to be user-friendly and is accompanied by documentation that guides the user through the process of network generation and sentiment propagation analysis.

An example of an FoF lexical network is given in Figure 1, showing the FoF_{wealth} network developed from the enTenTen corpus [64]. It can be observed that the central lexeme *wealth* attracts lexemes that are conceptually related. This network shows subcommunities that reflect different facets or contexts of the lexeme *wealth*. For example, lexemes such as *success*, *prosperity*, *health*, *happiness* and *abundance* form a conceptual subcommunity that refers to general well-being and fulfilment and is often associated with the achievement of a particular goal. In contrast, terms such as *income*, *status*, and *poverty* refer to financial status and reflect conditions or systems that influence *wealth*. This example not only illustrates the semantic proximity of the different lexemes to the source lexeme *wealth* but also highlights the different but interrelated conceptual domains it encompasses.

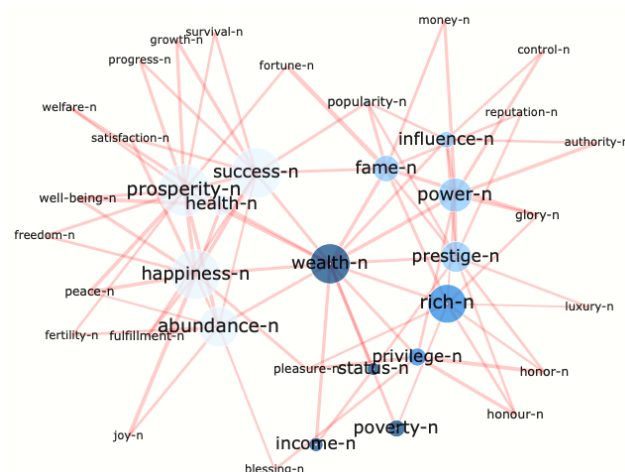


Figure 1. FoF_{wealth} lexical network of the noun lexeme *wealth*.

3.3. Assigned Dictionary Values of Lexemes

The work in [12] introduced a way to assign sentiment values to lexemes by computing the sentiment values of entire lexical networks. These lexical networks are constructed with lexemes as nodes and their dependency relations as edges. The structure of these weighted undirected graphs reveals which lexemes are represented by nodes with greater or lesser importance.

Sentiment estimation takes into account both the structure of the weighted undirected graph itself and the available dictionary sentiment values of the node lexemes. Typically, sentiment values are not known for all lexeme nodes in the lexical network, so the sentiment of a graph is calculated from the known sentiment values of the lexeme nodes in proportion to their centrality. In this way, the most important nodes in the network contribute the most to the sentiment assessment of the entire graph.

More precisely, given a lexical graph G with the set of nodes V and a centrality measure of node importance c , let $c(x)$ denote the centrality of a node x . Let $V_G^D \subseteq V$ denote the set of nodes $x \in V$ for which the numerical sentiment value of the node lexeme $x \in V$, $s(x)$ is known, i.e., appears in the dictionary D . The sentiment value of the graph G is computed by

$$GSV(G) := \frac{\sum_{x \in V_G^D} s(x) \cdot c(x)}{\sum_{x \in V_G^D} c(x)}. \quad (1)$$

The calculated sentiment value of a lexical graph is then assigned to the seed lexeme around which the graph was constructed. That is, new dictionary values are assigned to lexeme based on the corresponding lexical graphs of the same-part-of-speech collocates and their available sentiment values from the dictionary, and they are calculated for each of the relevant dictionary categories. The *assigned dictionary value* (ADV) of a node a using dictionary sentiment values in the category C_j of the dictionary D , denoted by $ADV(a, j)$, is defined for a non-empty set V_a^D as follows [12]:

$$ADV(a, j) := \frac{\sum_{x \in V_a^D} v_j(x) \cdot b(x)}{\sum_{x \in V_a^D} b(x)}, \quad (2)$$

where $v_j(x)$ is the sentiment value of the lexeme x in the category C_j of the dictionary D , also called *Original Dictionary Value* (ODV), V_a^D is the set of nodes $x \in FoF_a$ in the coordination-type lexical graph FoF_a of the source lexeme a for which $v_j(x)$ appears in D , and $b(x)$ is the *betweenness* centrality measure of the node x in the FoF_a graph.

The above formula facilitates the assignment of a newly computed sentiment value to the seed-lexeme of each FoF graph, effectively leveraging the collective sentiment values of related collocates within the graph. Consequently, lexemes are assigned with dictionary values, which are reflective of the aggregate sentiment derived by their respective FoF syntactic–semantic graphs.

3.4. Sentiment Dictionary Propagation Algorithm

The construction of a comprehensive sentiment dictionary begins with the creation of an extensive lexicon of lexemes. Such an extensive lexicon of words to which sentiment values are to be assigned is obtained from a selected large language corpus. We use existing sentiment dictionaries as a reference for the propagation of sentiment values. Using the method described in Section 3.2, we construct coordination-based lexical networks to compute the assigned dictionary values of lexemes (as per Equation (2)) for each lexeme of the extracted lexicon for which the sentiment value is not given by an existing dictionary. Multiple sentiment values can be calculated for each lexeme to expand available sentiment dictionaries and their categories. In this way, we obtain a much larger, multidimensional dictionary coverage.

The sentiment dictionary to be created consists of a set of lexemes and their respective sentiment values over a number of sentiment categories. In the dictionary, we include words

with part-of-speech tags (lempos) that are typically considered in sentiment dictionaries, i.e., nouns, verbs, adjectives, and adverbs.

The propagation algorithm uses a number of resources and has several parameters. For the specification of the algorithm, we use an intuitive, descriptive meta-language and the notation given in Table 1.

Table 1. Notation used in the sentiment dictionary propagation algorithm.

Notation	Denotation
C_i	a selected corpus, for $i \in \{1, \dots, k_C\}$;
D_i	a selected sentiment dictionary, for $i \in \{1, \dots, k_D\}$;
s_i	a category from one of the selected sentiment dictionaries, for $i \in \{1, \dots, r\}$;
$Dict(c)$	the corresponding sentiment dictionary, of which c is a category;
S_{pos}	the set of part-of-speech tags suitable for sentiment assignment;
$Pos(a)$	the set of all part-of-speech tags $p \in S_{pos}$ of the lexeme a in the selected corpora;
L_C	the set of lexemes from selected corpora with part-of-speech tags in S_{pos} ;
L'_C	the set of lempos from selected corpora with part-of-speech tags in S_{pos} ;
L_i	the set of lexemes that have a sentiment score in category s_i , for $i \in \{1, \dots, r\}$;
L'_i	the set of lempos that have a sentiment score in category s_i together with their $Pos(a, A)$ tags, for $i \in \{1, \dots, r\}$;
L''_i	the set of lempos that have a sentiment score in category s_i , which do not appear in the selected corpora, for $i \in \{1, \dots, r\}$;
L_S	the set of lexemes of dictionary S together with their tags, i.e., lempos of S ;
M_{multi}^{pos}	the set of lexemes of S with multiple pos-tags;
$ODV(x, j)$	the original sentiment value of lemma x in category j of a pos-untagged dictionary;
$ODV(x-t, j)$	the original sentiment value of lempos $x-t$ in category j of a pos-tagged dictionary;
$N(FoFa)$	the set of node lexemes of $FoFa$ graph;
$k(FoFa, j)$	the number of nodes in the $FoFa$ graph for which the sentiment value in category j of S is undefined, for $j \in \{1, \dots, r\}$;
$l(FoFa, j)$	the number of nodes in the $FoFa$ graph for which the sentiment value in category j of S is defined, for $j \in \{1, \dots, r\}$;
$z(FoFa, j)$	the proportion of nodes in the $FoFa$ graph for which the sentiment value in category j of S is defined, rounded to two decimal places, for $j \in \{1, \dots, r\}$;
$v_i(a)$	the sentiment value of lempos $a \in L$ in category s_i of S , for $i \in \{1, \dots, r\}$.

The proposed propagation algorithm itself is given below (Algorithm 1), while its high-level block diagram is shown in Figure 2. It provides an overview of how the sentiment dictionary is populated: first from input data extracted from selected corpora and selected sentiment dictionaries and then through an iterative process of propagating sentiment information.

Algorithm 1 Propagation of sentiment values

```

1: Select tagged corpora  $C_1, \dots, C_{k_C}$ 
2:  $C := C_1 \cup \dots \cup C_{k_C}$ 
3: Select sentiment dictionaries,  $D_1, \dots, D_{k_D}$ 
4: Form the list of all the corresponding dictionary categories,  $s_1, \dots, s_r$ 
5: for  $1 \leq i \leq r$  do set  $Dict(i) := A$ , where  $s_i$  is a category of the dictionary  $A$ ;
6: end for
7:  $S_{pos} := \{n, v, adj, adv\}$ ;
8:  $L_C := \{x \mid \text{lemma } x \text{ appears in } C \text{ with tag } t \in S_{pos}\}$ 
9:  $L'_C := \{x-p \mid \text{lemma } x \text{ appears in } C \text{ with tag } p \in S_{pos}\}$ 
10: for  $x \in L_C$  do  $Pos(x) := \{t \mid x-t \in L'_C\}$ 
11: end for
12:  $M_{ulti}^{pos} := \{x \mid x \in L_C, |Pos(x)| > 1\}$ 
13: for  $1 \leq i \leq r$  do  $L_i := \{x \mid \text{sentiment score for lemma } x \text{ occurs in category } s_i\}$ 
14:  $L'_i := \{x-p_x \mid x \in L_i, p_x \in Pos(x)\}$ 
15: end for
16:  $L_S := L'_C \cup L'_1 \cup \dots \cup L'_r$ 
17: for  $1 \leq i \leq r$  do
18:   for each  $a \in L_S$  do
19:     if  $a \in L'_i$  then  $v_i(a) := ODV(a, s_i)$ 
20:     else  $v_i(a) := undef$ 
21:     end if
22:   end for
23: end for
24: Select parameters for the construction of FoF lexical networks,  $n$  and  $m$ 
25:  $j := 0$ 
26:  $j := j + 1$ 
27: if  $j > r$  then STOP
28: end if
29:  $X := \{x-t \in L_S \mid v_j(x-t) = undef, t \in Pos(x)\}$ 
30: for each  $a \in X$  do build a  $n \times m$  FoF $_a$  network
31: end for
32: for each  $a \in X$  do
33:    $k(FoF_a, j) := |\{x \in N(FoF_a) \mid v_j(x) \neq undef\}|$ 
34:    $l(FoF_a, j) := |\{x \in N(FoF_a) \mid v_j(x) = undef\}|$ 
35:    $z(FoF_a, j) := \text{round}(\frac{k(FoF_a, j)}{l(FoF_a, j) + k(FoF_a, j)}, 2)$ 
36: end for
37:  $m_{stop} := \max\{k(FoF_a, j) \mid a \in X\}$ 
38: if  $m_{stop} = 0$  then GOTO 26
39: end if
40:  $m_{ax} := \max\{z(FoF_a, j) \mid a \in X\}$ 
41:  $H := \{a \in X \mid z(FoF_a, j) = m_{ax}\}$ 
42: for each  $x-t \in H$  do
43:   if  $Dict(j)$  is pos-tagged then  $v_j(x-t) := ADV(x-t, j)$ 
44:   else [ $Dict(j)$  is pos-untagged]
45:     if  $x \in L_j$  then
46:       if  $x \notin M_{ulti}^{pos}$  then  $v_j(x-t) := ODV(x, j)$ 
47:       else [ $x \in M_{ulti}^{pos}$ ]  $v_j(x-t) := \frac{1}{2}(ADV(x-t, j) + ODV(x, j))$ 
48:       end if
49:     else [ $x \notin L_j$ ]  $v_j(x-t) := ADV(x-t, j)$ 
50:     end if
51:   end if
52: end for
53:  $X := \{x-t \in L_S \mid v_j(x-t) = undef, t \in Pos(x)\}$ 
54: GOTO 32

```

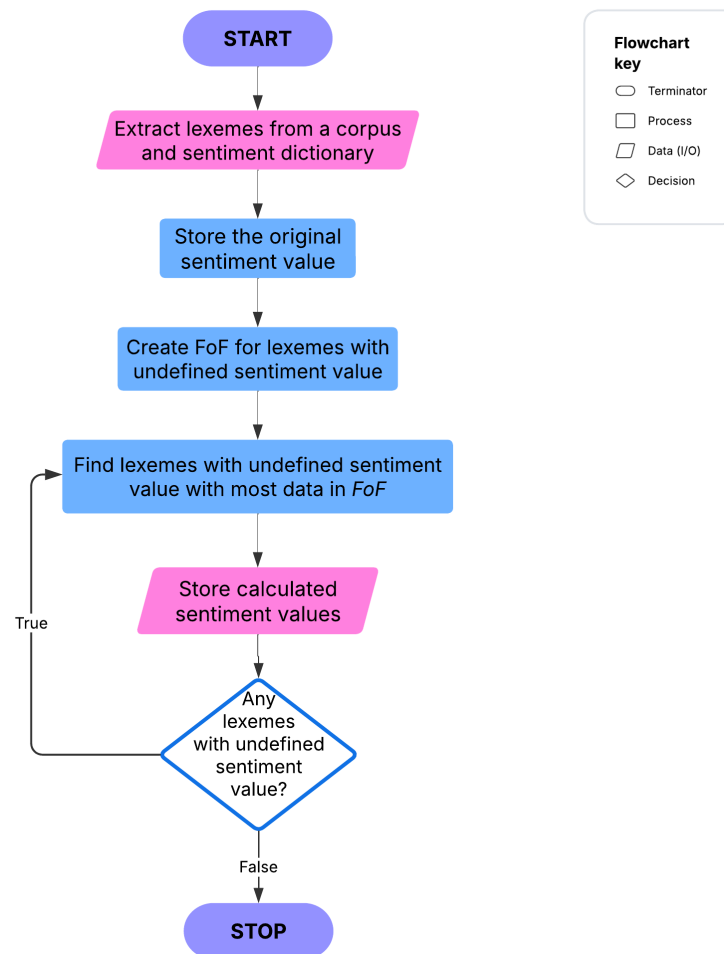


Figure 2. High-level block diagram of the sentiment dictionary propagation algorithm.

The procedure starts with the selection of corpora together with one or more sentiment dictionaries and their corresponding categories (steps 1–6). The list of dictionary entries of S , L_S is extracted from the corpora and the selected dictionaries (steps 7–16). The dictionary entries are lexemes extracted from the selected corpora and sentiment dictionaries together with the relevant part-of-speech tags, i.e., nouns, verbs, adjectives and adverbs, such as the lempos *power-n*, (*to*) *power-v*. Lexemes extracted from pos-untagged sentiment dictionaries are also tagged with relevant, possibly multiple part-of-speech tags.

Note that the original sentiment score of a lexeme in a pos-untagged sentiment dictionary is considered cumulatively for all its different pos-tags.

Initially, the original sentiment values are copied from all selected dictionary categories for the lexemes for which these values already exist for a specific part of speech in the corresponding pos-tagged dictionary (step 19). Otherwise, the initial sentiment value of lexemes in S are set as undefined (step 20). The ODV scores from pos-untagged sentiment dictionaries are not simply carried over for lexemes with multiple part-of-speech tags but are instead propagated, taking into account both the original ODV value and the ADV value for a particular part of speech.

The next phase of the algorithm is the projection of sentiment values for lexemes and categories for which the original value is not available in the corresponding dictionary. The values for these remaining lexemes in the lexicon S are then filled in by category. This is accomplished using the counter j (initialised and controlled in steps 25–27).

For each of the source lexemes with an undefined sentiment value in the category, we construct the $n \times m$ lexical FoF network using the ConGraCnet approach (step 30). This lexical network contains the highest collocates of the source lexeme a based on the

same part-of-speech [lexeme_1 and/or lexeme_2] coordination relation obtained using the parameters chosen in step 24.

As default values for these parameters, we take $n = 15$ and $m = 5$, i.e., set the default dimensions of the lexical FoF network to 15×5 .

Lexemes for which the sentiment value in the category is still undetermined are selected (formation of the set X , step 29). Among these, those lexemes are identified that contain the most information about the sentiment values in their FoF networks. That is, for each lexeme in $a \in X$, we determine $k(FoF_a, j)$ and $l(FoF_a, j)$, including the number of nodes and the proportion of nodes in FoF_a for which the sentiment value v_j is already defined in S (steps 32–34). The maximum value of the proportion of nodes with the defined value in the FoF_a graph, m_{ax} , is calculated (step 40). This value is then used to select the subset H of X that contains lexemes with the highest proportion of nodes in the FoF_a graph with an already assigned sentiment value (steps 33–41). The value m_{stop} (calculated in step 37) is used to check when all sentiment values in the current category have been assigned so that propagation proceeds to the next category.

Next, we assign a sentiment score to each lempo $x-t \in H$ since, by definition of H , the sentiment value of x is undefined for part of speech t .

That is, for each lempo $x-t$ in H , we check whether the sentiment category j corresponds to a pos-tagged or a pos-untagged sentiment dictionary. If $Dict(s_j)$ is pos-tagged, the sentiment value for $x-t$ is calculated as ADV for that part of speech (43). In case $Dict(s_j)$ is pos-untagged, we distinguish between several cases. In the case that the lexeme x with its $ODV(x)$ occurs in the corresponding untagged dictionary $Dict(s_j)$, we take this original sentiment score into account:

$$v_j(x-t) := \begin{cases} ODV(x, j) & \text{if } x \in L_j, x \notin M_{ulti}^{pos} \\ \frac{1}{2}(ADV(x-t, j) + ODV(x, j)) & \text{if } x \in L_j, x \in M_{ulti}^{pos} \\ ADV(x-t, j) & \text{if } x \notin L_j, \end{cases}$$

For lexemes with existing ODV values that refer to words with a single pos in the corpus, the ODV value from the categories of the pos-untagged sentiment dictionary is assigned to the lexemes with this identified part-of-speech tag (step 46). For example, the ODV of *elephant* is carried over to the lempo *elephant-n*. For lexemes with existing ODV values that refer to words with multiple pos in the corpus, the sentiment value of lempo $x-t \in H$ is calculated as the mean of the ODV value of the lexeme x and the ADV value determined using FoF_a constructed over t part-of-speech collocations (step 47). For example, the ODV of the lexeme *power* is considered together with the ADV of the lempo *power-n* and the lempo *power-v* to propagate the respective noun and verb sentiment values. In the remaining case, if the lexeme x does not occur in the pos-untagged dictionary $Dict(s_j)$, there is no original sentiment score in the dictionary, so we set the propagated value for x and the part of speech t to $ADV(x-t, j)$ (step 49).

We iterate (step 54) the ranking of the lexemes to be assigned the sentiment value (recalculated in step 53) according to the number of nodes with available score, and we gradually assign the sentiment value to the lexemes with the most information about sentiment values in their FoF network (steps 32–38).

It should be emphasised that the more lexical nodes with sentiment values are present in the sentiment dictionary, the better the evaluation of the assigned sentiment value of the seed lexeme. Consequently, our approach proposes to first compute the sentiment score for the candidate lexemes with the most sentiment score information in their lexical FoF networks and gradually populate the dictionary. This is an important feature of our propagation algorithm, as it enables the iterative projection of sentiment load from the

original dictionary from the conceptual domains with the most information and gradually transfers the affective load to less covered conceptual domains.

4. Propagating the Sentiment-hr Dictionary from Coordination Based Lexical Graph

As described in Section 3, our sentiment propagation methodology is designed to improve sentiment dictionaries in any language provided that a lexical graph capturing semantic similarity is available. To test the effectiveness of this approach, we have developed Sentiment-hr, a sentiment dictionary for Croatian, which is the most comprehensive, publicly available sentiment dictionary for the Croatian language.

The first step in this process was to create a syntactic-semantic embedding graph using the large Croatian web corpus hrWaC 2.2, which contains 1.4 Giga words. According to our ConGraCNet methodology, we created the lexical graph structure required for the subsequent propagation of sentiment values [23]. The lexical graph includes 990,327 lexical nodes divided into 566,961 nouns, 188,298 adjectives, 27,014 adverbs and 208,054 verbs. It has 11,778,373 edges, which indicates extensive lexical connections. The analysis of edges shows a broad lexical scope of 6,660,582 edges connecting nouns, while adjectives have 1,503,228 edges, adverbs 196,167 edges and verbs 3,418,396 edges. The granularity of these connections forms a crucial basis for sophisticated sentiment propagation that goes beyond mere lexical lists and leads to a more integrated semantic network.

In addition to the unique lempos (lemma + pos) value, the node attributes in this lexical graph also contain the frequency of their occurrence in the corpus. The edge attributes, on the other hand, contain the frequency of co-occurrence and the logDice score—a statistical metric that assesses the strength of association between two lexemes based on their observed co-occurrence compared to what would be expected if they were independent, normalised by their frequency sum. This measure is particularly sensitive to significant but less frequent co-occurrences and is therefore ideal for highlighting meaningful lexical relationships within the corpus. These node and edge attributes are an essential part of the construction of a lexical embedding graph, which is subsequently used as a substrate for the construction of lexical FoF networks, providing a structured approach to semantic network analysis.

This graph, encapsulated as an igraph object [65], provides a robust structure for the application of our sentiment propagation algorithm. In addition, the full graph can be downloaded as igraph object [23], allowing wider access and use in sentiment analysis research and application.

After the construction of the network is completed, the sentiment values from the SenticNet sentiment dictionary are integrated. For our project, we chose SenticNet 6 (2020) as the cornerstone of the ConGraCNet application because it has an easily accessible API, and the Python library senticnet is able to efficiently retrieve sentiment dimensions and values. SenticNet 6 is known for its ability to reconcile symbolic and subsymbolic approaches. It masters the challenges of sentiment analysis by combining traditional linguistic methods with modern computational strategies [10].

The choice of SenticNet 6 was not only technical but also strategic, as it is adaptable and offers comprehensive language support, which is crucial for the project's multilingual ambitions. This choice facilitates the continuous development of our sentiment analysis framework and enables the integration of future SenticNet 6 iterations and enhancements.

This integration represents a strategic decision to use a dynamic and comprehensive sentiment analysis tool that underpins our goal of developing a comprehensive sentiment dictionary. By using SenticNet 6, we benefit from its extensive sentiment lexicon. The adaptability of our methodology also ensures that the approach exemplified by the ex-

tension of the Croatian sentiment dictionary is applicable to other languages and sentiment dictionaries utilising the BabelSenticNet resources [26].

The described procedure resulted in the assignment of initial sentiment values from SenticNet 6 with a script to match lemmas nodes from a SenticNet 6 word list. The initial SenticNet 6 values were assigned to 11,726 nodes within the lexical graph, which were systematically distributed across different word types: 6121 nouns, 2533 adjectives, 1450 adverbs and 1622 verbs. Since there is no POS in SenticNet, the same multiPOS words have the same assigned values for different POS. The nodes in a lexical graph were labelled with SenticNet categories: ‘polarity value’, ‘pleasantness’, ‘attention’, ‘sensitivity’, and ‘aptitude’, adding the expression of the original value ‘sentic_odv’ and the measure of assigned value certainty ‘adv_cert’. This step created a basic sentiment framework throughout the network, which facilitated the subsequent propagation of sentiment values throughout the network.

After setting up a basic sentiment framework within the network, we implemented the sentiment propagation Algorithm 1 to fill unassigned lexical nodes, as described in Section 3.4. Parameters used for this implementation include the following:

- Corpus: hrWaC [66];
- Sentiment dictionary: SenticNet 6;
- lexical FoF network parameters: $n = 15$, $m = 5$ resulting in a 15×5 FoF network.

The sentiment values were propagated throughout the network and resulted in a graph with 953,482 lexical nodes with computed sentiment values, including 536,079 nouns, 185,591 adjectives, 25,529 adverbs and 206,283 verbs, as shown in Table 2. A total of 25,119 nodes remained without calculated sentiment values due to inadequate lemmatisation, including 24,761 nouns, 174 adjectives, 35 adverbs, and 149 verbs. These lexemes remain unconnected in the network and hinder the propagation of sentiment.

Table 2. Summary of Sentiment-hr analysis metrics.

Metric	Count
Number of nodes	990,327
Number of edges	11,778,373
Number of nodes with sentic values	11,726
Nouns	6121
Adjectives	2533
Adverbs	1450
Verbs	1622
Number of nodes with calculated values	953,482
Nouns	536,079
Adjectives	185,591
Adverbs	25,529
Verbs	206,283
Number of nodes without calculated values	25,119
Nouns	24,761
Adjectives	174
Adverbs	35
Verbs	149

The proposed methodology underlines the importance of precise lemmatisation in the creation of effective corpus-based lexical networks and demonstrates the potential of syntactic–semantic networks to facilitate the propagation of sentiment values on a broad scale, thus significantly improving the coverage and dimensionality of sentiment dictionaries for languages such as Croatian.

The resulting resource is the most comprehensive sentiment dictionary in Croatian with coverage about five times larger than the original dictionary, which was the Croatian part of SenticNet 6. It is a freely accessible resource [25].

5. Extracting Sentiment Values from Large Language Models

In an era dominated by the exponential generation of data, researchers are constantly looking for efficient and robust tools to analyse and explore vast datasets. Rapid advances in artificial intelligence technologies have simplified this endeavour, especially when innovative and powerful AI models are used. The practice of comparing emerging AI methods with traditional non-AI approaches has become an integral part of contemporary research. In the context of our study, the comparison between conventional methods and state-of-the-art AI models is fruitful as it provides valuable insights, validates different viewpoints and sheds light on the significant advances in the technological landscape.

In the field of natural language processing (NLP), a category of sophisticated algorithms known as large language models (LLMs) has been developed. These models are the result of extensive training on huge corpora of text data, enabling them to recognise and interpret patterns and thus approach a form of “understanding” that enables interactions similar to human engagement. LLMs have a remarkable ability to generate texts on specific topics, respond to queries, produce images and music that embody certain characteristics, and—crucially for the purpose of our study—assess the sentiment contained in words, sentences or whole texts [67].

5.1. Sentiment-hr AI Dictionary Propagation

In our endeavour to create a sentiment dictionary, we tried an alternative approach using OpenAI’s large language model GPT-4, which is based on the large training data and the neural network transformer architecture [68] and is known for its effectiveness on sequential inputs. The extraction method involves programmatically analysing the sentiment of lexemes through a script that interacts with the GPT-4 API. The script processes lexemes in batches of five and queries the GPT model for sentiment analysis. The script has robust error handling and state management to ensure continuity and efficiency in processing.

The core of this process is the system query and the user query, which was created for GPT-4 with a simple instruction:

```
Propose sentiment value for a lexeme, presented as lempos, with lemma as the
first part and part-of-speech as the last part.
Lempos are lexical concepts in {language} language. Write the sentiment po-
larity and pleasantness as a fine-grained float value in a range from -1.00
to 1.00.
For each lempos write a response only in JSON format with keys:
lemma, part_of_speech sentiment_polarity, pleasantness.
Output JSON results as dictionaries, separated with commas, do not make a list
out of multiple dictionaries.
Target lempos.
```

The system prompt sets the context for the AI by describing it as a linguist capable of assigning sentiment values to lexical items in multiple languages. This feature enables the extraction of sentiment values across a broad linguistic spectrum and makes the approach adaptable for the creation of sentiment lexicons in different languages. In the user query, GPT-4 was tasked with assigning sentiment values to the first 49,338 lexemes extracted from the Sentiment-dictionary-hr list, which is sorted by degree. Responses are expected in a structured JSON format containing both sentiment polarity and pleasantness scores for each lexeme.

The GPT model generated grained sentiment values in categories: polarity value and pleasantness, which were both in the interval $[-1, 1]$ for 49,338 lexemes from the graph-based lexicon. The produced sentiment dictionary resource Sentiment-hr-AI is available online [27].

It is interesting to compare the sentiment values of Sentiment-hr-AI with Sentiment-hr. A comparison of some sentiment values for the graph-based propagation algorithm presented in Section 3.4 and using GPT-4 is shown in Table 3.

Table 3. Comparison of sentiment values.

Lempos	Graph-Based Propagation Algorithm	GPT-4
<i>dio-n</i> (<i>part-n</i>)	0.073946517	0
<i>ruka-n</i> (<i>hand-n</i>)	0.136528296	0.1
<i>aktivnost-n</i> (<i>activity-n</i>)	0.266668523	0.2
<i>gubitak-n</i> (<i>loss-n</i>)	−0.271211855	−0.6
<i>volja-n</i> (<i>will-n</i>)	0.239851995	0.6
<i>komentar-n</i> (<i>comment-n</i>)	0.231013698	−0.2
<i>obveza-n</i> (<i>obligation-n</i>)	0.492960287	−0.2
<i>ekran-n</i> (<i>screen-n</i>)	−0.300003924	0.1

These particular lexemes were selected on the basis of the comparison of their sentiment values obtained by GPT-4 queries and by the graph-based propagation algorithm presented in Section 3.4. While the first three lexemes show only minimal differences in the sentiment values, considerable discrepancies in sentiment intensity can be observed for *gubitak-n* (loss-n) and *volja-n* (will-n). The last three lexemes in Table 3 show such large differences that they are assigned opposite sentiment polarities—positive in one model and negative in the other. The sentiment values determined with the graph-based propagation algorithm can be easily checked, and their derivation can be systematically traced. In contrast, the values generated by the LLM are inherently opaque, as their calculation is based on extensive data from the internet.

Although both methods represent a significant quantitative progress, we make no claim to the qualitative assessment or comparison of sentiment values across dictionaries. Indeed, there is no objective reference value that would serve as a gold standard that could be achieved. By definition, sentiment dictionaries are subjective in nature. The data they contain depend on the broad cultural environment in which it was created and is therefore inherently characterised by variability due to language, context and subjectivity.

5.2. Analysis of LLM Sentiment Extraction

Utilising the deep learning capabilities of large-scale language models such as GPT-4 offers significant advantages, including the ability to process lexemes in multiple languages and the flexibility to adapt to different linguistic nuances and sentiment dimensions. However, this approach is not without its drawbacks. The “black box” nature of GPT-4’s processing and decision-making mechanisms means that the reasons for the sentiment values assigned are not transparent, and that presents a challenge for validation and cus-

tomisation. In addition, the reliance on a third-party API leads to potential access, cost and privacy issues.

Despite these challenges, the use of GPT-4 to extend the sentiment lexicon represents an exciting innovation in sentiment analysis. It complements traditional methods by providing an efficient, scalable and language-independent sentiment value mapping tool that opens up new possibilities for the exploration and application of sentiment analysis.

When analysing the sentiment extraction with GPT-4 from OpenAI, a scatterplot was created in Figure 3 to illustrate the relationship between two primary dimensions: sentiment polarity and pleasantness. Sentiment polarity, shown on the x -axis, is a numerical representation of the emotional charge of a lexeme, ranging from negative -1.00 to positive $+1.00$. Pleasantness, shown on the y -axis, quantifies the degree of positive emotional content associated with a term. Again, the scale ranges from -1.00 to 1.00 with higher values indicating a more positive emotional response. This two-dimensional analysis aims to shed light on the complicated relationship between the emotional charge of words and their ability to evoke a positive affective state.

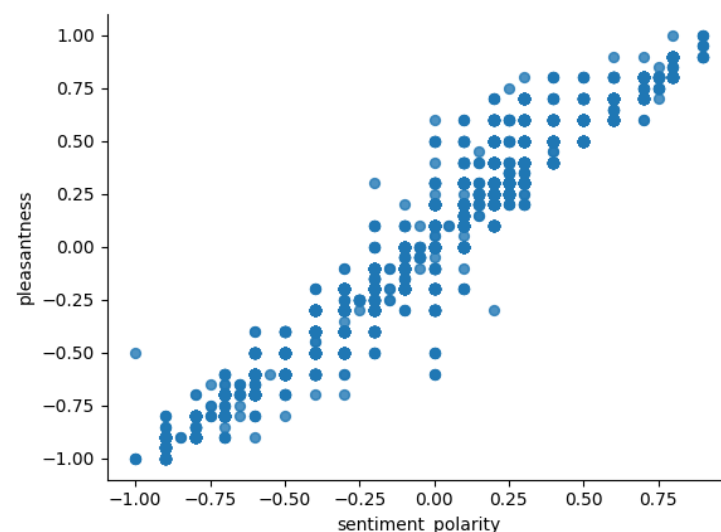


Figure 3. Scatterplot depicting sentiment polarity compared to pleasantness, extracted from GPT-4.

Looking at the scatterplot diagram, a trend can be recognised: when the polarity of feelings increases toward the positive end of the spectrum, pleasantness increases at the same time. This trend indicates a correlation between the two measures and suggests that words with a positive sentiment are generally associated with a higher level of pleasantness. The density of data points forming an ascending diagonal path from the lower left quadrant (negative sentiment polarity and unpleasantness) to the upper right quadrant (positive sentiment polarity and pleasantness) visually represents this relationship.

However, it is noticeable that there is a significant aggregation of data points centred around the neutral zero mark on the axis of sentiment polarity with a wide range of pleasantness scores. This clustering suggests that certain lexemes maintain neutral sentiment polarity but still exhibit varying degrees of pleasantness. This phenomenon emphasises the non-linear and complex nature of the relationship between sentiment polarity and pleasantness. It suggests that words can have a neutral sentiment but still have different emotional weight in terms of pleasantness.

The graphical representation of sentiment polarity and pleasantness extracted from GPT-4 provides valuable insight into the model's nuanced understanding of language. It reflects the model's ability to not only categorise words along a positive–negative spectrum but also to recognise the subtleties of the emotional connotations of words. This analysis

demonstrates GPT-4's advanced ability to recognise the gradient and complexity of emotions that language can express, which is crucial for the development of more sophisticated sentiment analysis tools.

In our comparative analysis, we evaluated the distribution of sentiment polarity values derived from two different methods: hrWac and SenticNet 6 graph-based propagation with sentiment values extracted from the large language model GPT-4. As shown in Figure 4, the distribution generated by the graph-based propagation method is centred on the neutrality axis and resembles a normal distribution, indicating a balanced assimilation of positive and negative sentiment across the lexicon. This symmetry suggests a nuanced representation of sentiment polarity, with the extremes tapering off, which may indicate a comprehensive sentiment landscape captured by the graph-based method.

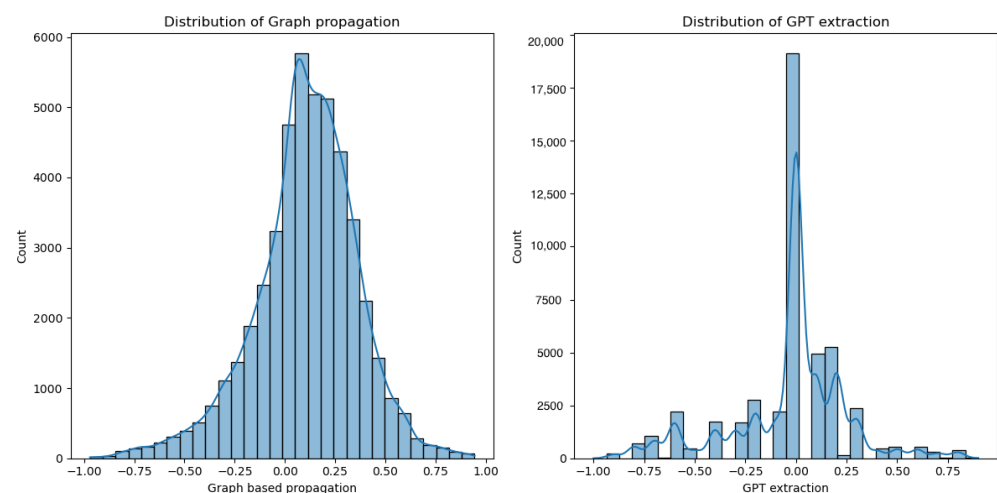


Figure 4. Comparison of sentiment polarity distributions: Sentiment-hr (Left) vs. Sentiment-hr-AI (Right).

In contrast, the GPT-based extraction showed a different distribution. The polarity values are predominantly grouped around the neutrality point with a pronounced peak at zero. This concentration of neutral sentiment values indicates a possible tendency of the GPT model towards non-committal sentiment categorisation. Furthermore, the presence of multiple local peaks at different polarity levels suggests that the GPT model tends to categorise sentiments in discrete clusters, which was possibly due to the fact that it was trained on a diverse corpus of delimited sentiment expressions.

The divergence in the distributions shows the different capabilities of the individual methods for analysing sentiment. The uniform distribution of the graph-based propagation method is suitable for applications that require subtle gradations of sentiment. In contrast, the segmented distribution profile of the GPT model may be suitable for contexts where clear categorisation of sentiment is of paramount importance.

6. Discussion

We have approached the problem of limited resources in sentiment analysis in two ways. The large language model gives us the ability to choose many factors, including the version of the model, the system prompt, and the parameters for a more dynamic and interactive connection with the language model. However, the results obtained with GPTs, especially in the context of language generation or comprehension tasks, can be difficult to explain as the specific methods, parameters and training data used to develop the model are not fully disclosed or transparent. Since GPTs are complex neural network architectures, it is difficult to explain step-by-step every output they produce. Our traditional algorithmic approach gives us an algorithm that is explainable and easy to understand and can be

broken down into steps that we can examine and, if necessary, change to improve the algorithm. Everything that has been said here contributes significantly to making this traditional scientific approach very standardised and understandable.

In the following subsection, we address various factors that influence the propagation and compare the results obtained.

6.1. Sentiment Dictionary Selection

Our approach is inherently versatile and suitable for improving sentiment dictionaries that characterise the emotional properties of lexemes by interval-based numerical values. This numerical representation enables the precise quantification of sentiment and facilitates the fine-grained analysis required for complex sentiment mapping tasks. The numerical intervals represent a spectrum of affective qualities ranging from strongly negative to strongly positive, allowing a nuanced distinction of emotional valence and intensity.

The current framework can potentially be adapted to include sentiment classifications with discrete numerical values or even categorical classes. Such adaptations would make it possible to create emotion dictionaries that classify emotions in a more segmented way, perhaps reflecting different emotional states or intensities without relying on a continuous scale. For example, a lexeme could be categorised into discrete emotion classes, such as ‘slightly positive’, ‘moderately positive’ or ‘very positive’, with each class corresponding to a numerical range or discrete value in the dictionary.

The ability to work with categorical sentiment values would open up avenues for sentiment analysis where cultural or contextual nuances are better expressed in categories rather than numerical gradations. For example, the sentiment associated with a word such as ‘victory’ might be seen as uniformly positive in different contexts and therefore better suited to a categorical approach rather than a nuanced numerical gradation.

However, extending our approach to such classifications leads to potential challenges, as categorical sentiment assignments require a different methodological treatment than interval-based numerical representations. Categorical classifications may not correlate linearly with the intensity of emotions and may involve more subjective interpretations. They also require additional processing to determine thresholds or rules for classification, which may differ significantly across languages and cultural contexts.

While our work currently focuses on interval-based numerical sentiment scores, the prospect of refining the model to integrate discrete and categorical sentiment classifications is an interesting direction for future work. This extension would improve the versatility of sentiment dictionaries and make them more adaptable and applicable to a wider range of linguistic and analytical scenarios. In this way, future research could provide an even more comprehensive set of tools for sentiment analysis, covering a wider range of emotional expressions within and across different linguistic landscapes.

6.2. Lexical Graph Model Selection

The versatility of the propagation Algorithm 1 presented in Section 3.4 is one of its most outstanding features. It offers users the autonomy to craft a lexical graph model that meets their specific analytical requirements. At the centre of such a lexical graph model is the concept of representing lexemes as nodes in a network with the connections between these nodes being established on the basis of synonymic relations. These relationships are quantified by weights, or the strength of association with the root node, effectively mapping the semantic proximity between the lexemes.

Such synonym-like relationships are an indicator of semantic similarity or relatedness, which, if quantified, can significantly improve the richness of the sentiment dictionary. The quantitative representation of the weights assigned to these relationships is an indicator

of the depth of sentiment association. A higher weight indicates a stronger synonymy or a deeper sentiment association, which in turn indicates a greater likelihood that the sentiment value of the root node is also shared by the associated lexeme.

Various data sources can be used to create this graph model, each of which has its own advantages. For example, corpus-based coordination collocates provide a data-driven foundation that is rooted in actual language usage and captures the natural coordination patterns of words. Synonym lexicons, on the other hand, offer a curated inventory of semantic relationships, which are usually derived from linguistic research and science. Word embeddings derived from machine learning algorithms encapsulate semantic relationships in a high-dimensional space, capturing nuances that may not be immediately recognisable from a superficial analysis. Similarly, large language models serve as a powerful tool for extracting complex lexical relations by drawing on the extensive knowledge encoded in their parameters.

When these elements are summarised in a lexical graph, the sentiment values of lexemes cannot be calculated in isolation but rather in the context of a rich network of semantic relationships. This contextuality is crucial for the creation of a sentiment dictionary that accurately reflects the different dimensions of sentiment as expressed and understood in natural language. This flexibility is not only a technical advantage but a methodological enrichment that adds contextual and cultural specificity to sentiment dictionaries.

A lexical model embodies a network of abstract conceptualisations that have crystallised within a particular community. These patterns are shaped by a variety of factors, including cultural norms, domain-specific nuances and even the temporal shifts that languages undergo over time [69]. Such a model goes beyond a static representation of lexical meanings. It is a dynamic representation of how a community understands and uses language. The choice of lexical graph model therefore has a major influence on the result of the sentiment analysis. For example, a model created with lexical relations from classical literary works gives a sentiment dictionary a very different cultural essence than a model derived from contemporary discourse in social media.

The emotional charge inherent in the lexemes from the selected sentiment dictionaries is then transferred to other lexemes to ensure that the transferred sentiment is coherent with the sense structures and emotional undercurrents specific to the lexical model used. For example, the lexeme *corona* may take on different affective properties when analysed against the background of historical literature than in a modern lexical model derived from texts about the pandemic.

In essence, the parameterised design of the algorithm facilitates the creation of sentiment dictionaries that not only serve sentiment analysis but also become repositories of cultural and linguistic intelligence capable of reflecting the rich diversity of human emotions as expressed through language.

6.3. On Lexical FoF Network Parameters

The construction of the lexical FoF networks follows a systematic approach in which each network starts from a source or seed lexeme. For each source lexeme, we identify a number of directly connected friend nodes that represent the first-degree connections in the network. Then, for each of these nodes, we identify additional friends that represent second-degree connections to the lexical FoF network. The result of this method is a richly connected FoF network that captures both direct and extended lexical relations and provides a comprehensive overview of the semantic field around each source lexeme.

When constructing the ConGraCNet FoF collocation networks, we strategically opted for standard parameters, namely $n = 15$ and $m = 5$, to delimit the dimensions of the FoF network as 15×5 . This choice of parameters is not arbitrary but is guided by the goal of

achieving a comprehensive coverage of lexemic associations without diluting the quality of the network with superfluous nodes. The justification for these parameter values lies in network theory and the principle of semantic proximity, which states that the affective evaluation of a lexeme is most strongly influenced by its closest lexical associations.

The selection of 15 direct collocates for each target lexeme ($n = 15$) ensures that the immediate semantic field is sufficiently represented and captures a set of lexemes that are most frequently and strongly associated with the target. By extending this network with five collocates for each of these lexemes ($m = 5$), we capture secondary but still significant semantic influences and enrich the network with depth and semantic domain context.

However, if you expand the network parameters indiscriminately to include more nodes, lexemes with marginal semantic relations could be included unintentionally. Such nodes contribute minimally to the assigned dictionary value (ADV) due to their peripheral relationship and lower centrality within the network. The inclusion of these nodes would likely lead to noise rather than valuable sentiment insight and potentially bias the sentiment analysis.

This careful calibration of FoF parameters is illustrated by examining sentiment propagation within lexemes associated with emotion-laden concepts such as ‘love’ or ‘grief’. Primary collocates of ‘love’ include, for example, ‘affection’, ‘passion’ and ‘romance’, which have a direct effect on the evaluation of sentiment. Secondary collocates such as ‘relationship’ or ‘heartbreak’ provide additional context that enriches the understanding of the feelings. Conversely, adding more distant collocates such as ‘date’ or ‘dinner’ may carry little emotional weight and therefore have little effect on the sentiment rating.

Therefore, the parameters we propose are designed to strike a balance between a comprehensive and an efficient network tailored to capture key sentiment information. This optimisation reflects a nuanced understanding of lexical semantics and network dynamics, which is crucial for the development of accurate and context-sensitive sentiment dictionaries.

6.4. Propagation Certainty: Proportion of Known to Overall Nodes in the Lexical FoF Network

The propagation certainty (ADV_CERT) is measured as the ratio of the number of nodes with known values to the total number of nodes in the lexical graph across computed lexical FoF networks. The ADV_CERT values propagated for the Sentiment_hrWac graph using the SenticNet 6 dictionary are shown in Figure 5 and summarised in Table 4. The diagram uses a logarithmic scale to better visualise the wide range of assignment reliability within the networks.

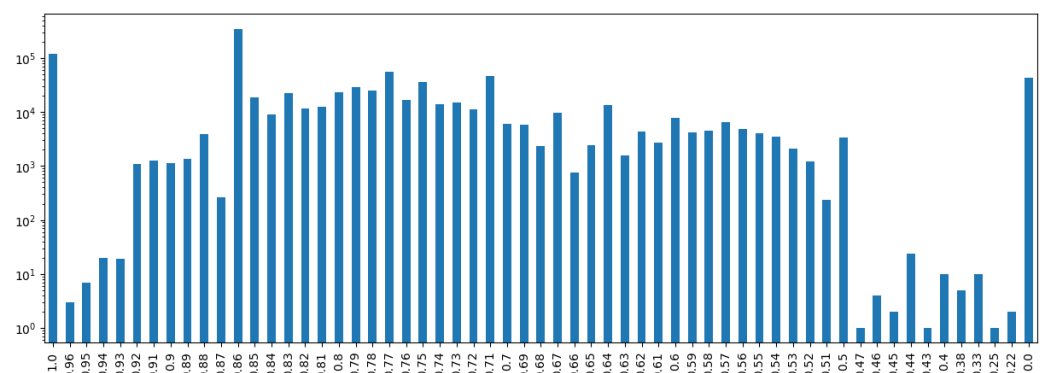


Figure 5. Distribution of the proportion of propagation certainty across 15×5 lexical FoF networks.

Table 4. ADV proportions certainty of nodes for which sentiment values have already been calculated in relation to the total number of nodes within each lexical FoF network.

ADV Proportion		Count	ADV Proportion		Count
0	1.00	118,134	30	0.67	9688
1	0.96	3	31	0.66	744
2	0.95	7	32	0.65	2385
3	0.94	20	33	0.64	13,440
4	0.93	19	34	0.63	1577
5	0.92	1081	35	0.62	4421
6	0.91	1247	36	0.61	2688
7	0.90	1146	37	0.60	7663
8	0.89	1358	38	0.59	4135
9	0.88	3864	39	0.58	4576
10	0.87	263	40	0.57	6497
11	0.86	346,032	41	0.56	4841
12	0.85	18,665	42	0.55	3980
13	0.84	8876	43	0.54	3548
14	0.83	22,724	44	0.53	2061
15	0.82	11,440	45	0.52	1195
16	0.81	12,502	46	0.51	232
17	0.80	23,219	47	0.50	3403
18	0.79	28,875	48	0.47	1
19	0.78	24,684	49	0.46	4
20	0.77	56,362	50	0.45	2
21	0.76	17,049	51	0.44	24
22	0.75	35,588	52	0.43	1
23	0.74	13,986	53	0.40	10
24	0.73	15,289	54	0.38	5
25	0.72	11,021	55	0.33	10
26	0.71	45,660	56	0.25	1
27	0.70	6020	57	0.22	2
28	0.69	5749	58	0.00	43,082
29	0.68	2383			

When analysing the distribution of assigned sentiment values within a 15×5 FoF network, the bar chart illustrates the variance of the certainty levels of these assignments, denoted *adv_cert*, versus the frequency of nodes with this proportion, denoted *count*. The *adv_cert* metric, ranging from 0.00 to 1.00, quantifies the confidence in the sentiment value assigned to each node with 1.00 representing absolute certainty. The results show a significant concentration of nodes (118,134) at the top of certainty (*adv_cert* = 1.00), indicating a robust ability of the algorithm to accurately determine sentiment values for a substantial portion of the network. In contrast, at the lower end of the spectrum (*adv_cert* = 0.00) there are 43,082 nodes, indicating cases where the algorithm was unable to determine sentiment with any degree of certainty.

A substantial peak at *adv_cert* = 0.86 with 346,032 nodes indicates a certain threshold in the sentiment mapping process, where a large number of nodes could be classified with high confidence although not absolutely. This indicates the effectiveness of the sentiment transfer algorithm, which utilises known sentiment values within the FoF network to derive sentiment for a large majority of nodes with high accuracy. The distribution across the other certainty levels is relatively sparse with the number of mentions decreasing as certainty drops from 1.00 to 0.50. This taper reflects the design of the algorithm, which favours a conservative approach where sentiment is only assigned when there is a high degree of certainty. The presence of nodes with lower *adv_cert* values (below 0.50) is minimised, meaning that the algorithm either achieves a high degree of certainty or settles on a state of indeterminacy (*adv_cert* = 0.00).

This distribution emphasises the effectiveness of the graph-based propagation algorithm in assigning sentiment values within the FoF network with a clear preference for assignments with high propagation certainty. The data suggest that while the algorithm performs excellently in identifying sentiment with high certainty for the majority of nodes, there is a non-trivial subset of the network where sentiment assignment occurs with lower certainty or not at all. Some of these lexemes may have been automatically taken from different corpora, which sometimes contain misspelled words or words that are very rare in the language. Consequently, such words have a very low certainty, as the algorithm does not have a high confidence in the assigned sentiment value (if any was determined). The original dictionary can be refined by removing such misspelled entities or entities with no meaning. These issues with propagation certainty also point to potential areas where the algorithm can be refined or where additional context- or network-based indicators need to be included to improve the overall confidence and accuracy of sentiment assignment.

6.5. Comparison of Approaches: Traditional vs. AI-Enhanced Sentiment Analysis

In the field of sentiment analysis, traditional algorithmic approaches and AI-supported methodology represent two different paths, each with its own advantages and limitations. The traditional algorithmic approach used in our study does not require the extensive learning process that characterises AI models. It works with a relatively limited dataset consisting of a selected corpus and an elementary sentiment dictionary. Despite this apparent simplicity, the reliability of traditional methods in deriving sentiment scores is remarkable. The focus of our method is on the seed lexeme around which all polysemous meanings are captured and analysed. This comprehensive consideration of the multiple meanings of a lexeme in different linguistic contexts ensures that sentiment scores are computed with a high degree of semantic awareness and contextual sensitivity.

Conversely, training a large language model (LLM) like GPT-4 is a massive endeavour that requires more than just access to open-source code, large datasets and powerful computational resources. It is a journey through a maze of often undocumented techniques that require constant vigilance to minimise the occurrence of unpredictable behaviour or anomalies. With this in mind, the recent LLM360 initiative represents a paradigm shift in the practice of model sharing within the AI community. It seeks to go beyond the traditional model of simply disseminating model weights and evaluation results. LLM360 is committed to providing a holistic overview of the LLM training process—including all intermediate checkpoints, all training data, comprehensive metrics and full source code—thereby advocating for greater transparency and comprehensibility [70].

One of the well-known problems with AI-based approaches is that they often tend to generalise sentiment assignments over lexemes with common morphological roots. As a result, terms that have a common root but differ in meaning may be incorrectly assigned the same sentiment polarity. For example, in some widely used sentiment dictionaries for English, the lexemes *starcraft* and *stardust*, and *fish* and *unselfishness*, are, respectively, assigned identical sentiment values simply because of their syntactic similarity, even though they have different semantic interpretations. Our ConGraCNet algorithm, on the other hand, does not rely on syntactic similarity but assigns sentiment values according to semantic similarity.

Our study contrasts these two approaches to sentiment analysis and emphasises the robustness of the traditional method in capturing nuanced lexical meanings against the backdrop of the evolving transparency and reproducibility of AI. This comparison serves as a guide for future methods of sentiment analysis. It suggests that an optimal sentiment analysis tool could well be a hybrid that combines the algorithmic precision of traditional approaches with the evolving, data-rich insights of AI models, as promoted by initiatives

such as LLM360. The interplay of these approaches has the potential to create a framework for sentiment analysis that is both robust and nuanced and can handle the complexity of human language and emotions.

7. Conclusions and Future Work

In this study, we have presented advances in sentiment analysis that bridge the gap of scarce computational resources in sentiment analysis. We applied a graph-based method to improve the accuracy and efficiency of sentiment analysis and presented the ConGraCNet Sentiment Propagation algorithm and the underlying Syntactic–Semantic–hrWac Embedding Graph. The algorithm leverages the strengths of both manual annotation and computational modelling to improve the resources available for sentiment analysis. We demonstrated its linguistic power by creating the most comprehensive sentiment dictionaries for Croatian to date, Sentiment-hr. In addition, we constructed Sentiment-hr-AI dictionary using LLMs.

Our work illustrates the successful integration of corpus-based and AI-driven approaches that improve our framework for sentiment analysis. This approach not only contributes to the burgeoning field of sentiment analysis in languages with scarce computational resources such as Croatian but also sets a precedent for the development of sentiment analysis tools for other languages with poor resources. This research sets the stage for expanding the capabilities of sentiment analysis and serves as a call to action for further innovations in computational linguistics to promote a deeper understanding of sentiment in human communication.

Our methods have proven to be effective in providing nuanced sentiment interpretations, although they need further refinement and validation. The obtained Croatian language example dictionaries illustrate the complexity and opportunities associated with extending sentiment analysis to languages with limited resources. The transition to AI-powered approaches in the creation of sentiment dictionaries not only emphasises the technological advancement in the field but also sets the stage for more comprehensive and effective sentiment analysis in a wider range of languages.

We would like to extend our methods to more languages and improve sentiment prediction by incorporating additional linguistic features and taking into account the dynamic nature of language use. Future efforts will focus on expanding the coverage of the sentiment dictionary, refining our algorithm to adapt to contextual changes, and incorporating multimodal data for a comprehensive understanding of sentiment. We plan to develop tools for domain-specific sentiment analysis and investigate the impact of cultural and temporal factors on sentiment.

Author Contributions: Conceptualisation, T.B.K., S.B.B. and B.P.; methodology, T.B.K. and B.P.; software, S.B.B. and B.P.; validation, T.B.K., S.B.B. and B.P.; formal analysis, T.B.K., S.B.B. and B.P.; investigation, T.B.K., S.B.B. and B.P.; resources, B.P.; data curation, S.B.B. and B.P.; writing—original draft preparation, T.B.K., S.B.B. and B.P.; writing—review and editing, T.B.K., S.B.B. and B.P.; visualisation, T.B.K., S.B.B. and B.P.; supervision, T.B.K., S.B.B. and B.P.; project administration, T.B.K.; funding acquisition, S.B.B. All authors have read and agreed to the published version of the manuscript.

Funding: This work has been supported in part by the Croatian Science Foundation under the project IP-2024-05-3882 and the University of Rijeka under the projects uniri-iskusni-prirod-23-150 and UNIRI-human-18-243.

Data Availability Statement: The original contributions presented in this study are included in the article. Further inquiries can be directed to the corresponding authors.

Conflicts of Interest: The authors declare no conflicts of interest.

References

- Liu, B. *Sentiment Analysis and Opinion Mining*; Springer Nature: Cham, Switzerland, 2022.
- Liu, B. Sentiment Analysis and Subjectivity. In *Handbook of Natural Language Processing*; Chapman and Hall/CRC: Oxfordshire, UK, 2010; pp. 627–666.
- Barrett, L.F. Valence Is a Basic Building Block of Emotional Life. *J. Res. Personal.* **2006**, *40*, 35–55. [CrossRef]
- Al-Qablan, T.A.; Mohd Noor, M.H.; Al-Betar, M.A.; Khader, A.T. A Survey on Sentiment Analysis and Its Applications. *Neural Comput. Appl.* **2023**, *35*, 21567–21601.
- Badaro, G.; Jundi, H.; Hajj, H.; El-Hajj, W. EmoWordNet: Automatic Expansion of Emotion Lexicon Using English WordNet. In Proceedings of the Seventh Joint Conference on Lexical and Computational Semantics, Orleans, LA, USA, 5–6 June 2018; pp. 86–93.
- Cui, J.; Wang, Z.; Ho, S.-B.; Cambria, E. Survey on Sentiment Analysis: Evolution of Research Methods and Topics. *Artif. Intell. Rev.* **2023**, *56*, 8469–8510.
- Khoo, C.S.; Johnkhan, S.B. Lexicon-Based Sentiment Analysis: Comparative Evaluation of Six Sentiment Lexicons. *J. Inf. Sci.* **2018**, *44*, 491–511.
- Taboada, M.; Brooke, J.; Tofiloski, M.; Voll, K.; Stede, M. Lexicon-based methods for sentiment analysis. *Comput. Linguist.* **2011**, *37*, 267–307.
- Villanes, A.; Healey, C.G. Domain-specific text dictionaries for text analytics. *Int. J. Data Sci. Anal.* **2023**, *15*, 105–118.
- Cambria, E.; Li, Y.; Xing, F.Z.; Poria, S.; Kwok, K. SenticNet 6: Ensemble Application of Symbolic and Subsymbolic AI for Sentiment Analysis. In Proceedings of the 29th ACM International Conference on Information & Knowledge Management, Virtual, 19–23 October 2020; pp. 105–114.
- Imbir, K.K.; Duda-Golańska, J.; Wielgopolan, A.; Sobieszek, A.; Pastwa, M.; Zygierewicz, J. The Role of Subjective Significance, Valence, and Arousal in the Explicit Processing of Emotion-Laden Words. *PeerJ* **2023**, *11*, e14583.
- Ban Kirigin, T.; Bujačić Babić, S.; Perak, B. Lexical Sense Labeling and Sentiment Potential Analysis Using Corpus-Based Dependency Graph. *Mathematics* **2021**, *9*, 1449. [CrossRef]
- Zhang, C.; Chen, Z.; Yan, D.; Cao, J.; Zhang, Q. Is a Single Embedding Sufficient? Resolving Polysemy of Words from the Perspective of Markov Decision Process. In *Database Systems for Advanced Applications, Proceedings of the International Conference on Database Systems for Advanced Applications, Tianjin, China, 17–20 April 2023*; Springer: Cham, Switzerland, 2023; pp. 555–571.
- Zhang, X.; Mao, R.; He, K.; Cambria, E. Neuro-Symbolic Sentiment Analysis with Dynamic Word Sense Disambiguation. In *Findings of the Association for Computational Linguistics: EMNLP 2023*; Association for Computational Linguistics: Singapore, 2023; pp. 8772–8783.
- Kaity, M.; Balakrishnan, V. Sentiment Lexicons and Non-English Languages: A Survey. *Knowl. Inf. Syst.* **2020**, *62*, 4445–4480.
- Mabokela, R.; Schlippe, T. A sentiment corpus for South African under-resourced languages in a multilingual context. In Proceedings of the 1st Annual Meeting of the ELRA/ISCA, Marseille, France, 20–25 June 2022; pp. 70–77.
- Koto, F.; Beck, T.; Talat, Z.; Gurevych, I.; Baldwin, T. Zero-shot sentiment analysis in low-resource languages using a multilingual sentiment lexicon. *arXiv* **2024**, arXiv:2402.02113.
- ConGraCNet Code. Available online: <https://github.com/bperak/ConGraCNet> (accessed on 20 March 2025).
- Perak, B.; Ban Kirigin, T. Construction Grammar Conceptual Network: Coordination-Based Graph Method for Semantic Association Analysis. *Nat. Lang. Eng.* **2023**, *29*, 584–614.
- Ban Kirigin, T.; Bujačić Babić, S.; Perak, B. Graph-Based Taxonomic Semantic Class Labeling. *Future Internet* **2022**, *14*, 383. [CrossRef]
- Ban Kirigin, T.; Bujačić Babić, S.; Perak, B. Semi-Local Integration Measure of Node Importance. *Mathematics* **2022**, *10*, 405. [CrossRef]
- ConGraCNet Application. Available online: <https://polinom.uniri.hr> (accessed on 20 March 2025).
- Lexical Embedding Graph. Available online: <https://drive.google.com/file/d/1HB-o4YQcvIog4q7d2WNUte7Rjoi05EpS/view?usp=sharing> (accessed on 20 March 2025).
- Ljubešić, N. hrWaC and slWaC: Compiling Web Corpora for Croatian and Slovene. In *Text, Speech and Dialogue, Proceedings of the 14th International Conference on Text, Speech and Dialogue, TSD 2011, Pilsen, Czech Republic, 1–5 September 2011*; Springer: Berlin, Germany, 2011; pp. 395–402.
- Sentiment-hr Dictionary. Available online: <https://drive.google.com/file/d/1NRVA-ZFkhkBSrMma8jhcl0LHeGLBMJAv/view?usp=sharing> (accessed on 27 March 2025).
- Vilares, D.; Peng, H.; Satapathy, R.; Cambria, E. abelSenticNet: A commonsense reasoning framework for multilingual sentiment analysis. In Proceedings of the 2018 IEEE Symposium Series on Computational Intelligence (SSCI), Bangalore, India, 18–21 November 2018; pp. 1292–1298.
- Sentiment-hr-AI Sentiment Dictionary. Available online: <https://drive.google.com/file/d/13tNteYu6QmIAPOYefui4SiD4YuytB4MV/view?usp=sharing> (accessed on 27 March 2025).

28. OpenAI. OpenAI: Artificial General Intelligence (AGI). 2023. Available online: <https://openai.com/blog/new-models-and-developer-products-announced-at-devday> (accessed on 20 March 2025).
29. Mohammad, S.M. Sentiment Analysis: Detecting Valence, Emotions, and Other Affectual States from Text. In *Emotion Measurement*; Elsevier: Amsterdam, The Netherlands, 2016; pp. 201–237.
30. Mohammad, S.; Bravo-Marquez, F.; Salameh, M.; Kiritchenko, S. SemEval-2018 Task 1: Affect in Tweets. In Proceedings of the 12th International Workshop on Semantic Evaluation (SemEval-2018), New Orleans, LA, USA, 5–6 June 2018; pp. 1–17.
31. Wilson, T.; Wiebe, J.; Hoffmann, P. Recognizing contextual polarity: An exploration of features for phrase-level sentiment analysis. *Comput. Linguist.* **2009**, *35*, 399–433.
32. Barel, G.; Tsur, O.; Vilenchik, D. Acquired TASTE: Multimodal Stance Detection with Textual and Structural Embeddings. *arXiv* **2024**, arXiv:2412.03681.
33. Lane, J.; Saint-Amand, H. Sentiment Analysis in Financial Texts. In Proceedings of the 28th International Conference on Computational Linguistics, Online, 8–13 December 2020; Association for Computational Linguistics: Barcelona, Spain, 2020; pp. 1–12.
34. Alamoodi, A.H.; Zaidan, B.B.; Zaidan, A.A.; Albahri, O.S.; Mohammed, K.I.; Malik, R.Q.; Almahdi, E.M.; Tareq, Z.; Albahri, A.S.; Chyad, M.A.; et al. A Systematic Review and Meta-Analysis of Artificial Intelligence (AI)-Based Breast Cancer Diagnosis and Prognosis Models in Medical Imaging Modalities: Deep Learning and Hybrid Learning Approaches. *Comput. Biol. Med.* **2021**, *132*, 104357. [[CrossRef](#)]
35. Pang, B.; Lee, L. Opinion Mining and Sentiment Analysis. *Found. Trends Inf. Retr.* **2008**, *2*, 1–135. [[CrossRef](#)]
36. Wankhade, M.; Rao, A.C.; Kulkarni, C. A survey on sentiment analysis methods, applications, and challenges. *Artif. Intell. Rev.* **2022**, *55*, 5731–5780.
37. Bordoloi, M.; Biswas, S.K. Sentiment Analysis: A Survey on Design Framework, Applications, and Future Scopes. *Artif. Intell. Rev.* **2023**, *56*, 12505–12560.
38. Esuli, A.; Sebastiani, F. SENTIWORDNET: A Publicly Available Lexical Resource for Opinion Mining. In Proceedings of the Fifth International Conference on Language Resources and Evaluation (LREC'06), Genoa, Italy, 22–28 May 2006; pp. 1–5. Available online: http://www.lrec-conf.org/proceedings/lrec2006/pdf/384_pdf.pdf (accessed on 20 March 2025).
39. Barnes, J. Sentiment and Emotion Classification in Low-Resource Settings. In Proceedings of the 13th Workshop on Computational Approaches to Subjectivity, Sentiment, & Social Media Analysis, Toronto, ON, Canada, 14 July 2023; pp. 290–304.
40. Buechel, S.; Rücker, S.; Hahn, U. Learning and Evaluating Emotion Lexicons for 91 Languages. *arXiv* **2020**, arXiv:2005.05672.
41. Cortis, K.; Freitas, A.; Daudert, T.; Huerlimann, M.; Zarrouk, M.; Handschuh, S.; Davis, B. SemEval-2017 Task 5: Fine-Grained Sentiment Analysis on Financial Microblogs and News. In Proceedings of the 11th International Workshop on Semantic Evaluation (SemEval-2017), Vancouver, BC, Canada, 3–4 August 2017; pp. 519–535.
42. Baccianella, S.; Esuli, A.; Sebastiani, F. SentiWordNet 3.0: An Enhanced Lexical Resource for Sentiment Analysis and Opinion Mining. In Proceedings of the Seventh International Conference on Language Resources and Evaluation (LREC'10), Valletta, Malta, 17–23 May 2010.
43. Miller, G.A. WordNet: A Lexical Database for English. *Commun. ACM* **1995**, *38*, 39–41.
44. Cambria, E.; Poria, S.; Hazarika, D.; Kwok, K. SenticNet 5: Discovering Conceptual Primitives for Sentiment Analysis by Means of Context Embeddings. *AAAI Conf. Artif. Intell.* **2018**, *32*, 1795–1802.
45. Ahmed, M.; Chen, Q.; Li, Z. Constructing Domain-Dependent Sentiment Dictionary for Sentiment Analysis. *Neural Comput. Appl.* **2020**, *32*, 14719–14732.
46. Cambria, E.; Fu, J.; Bisio, F.; Poria, S. AffectiveSpace 2: Enabling Affective Intuition for Concept-Level Sentiment Analysis. *AAAI Conf. Artif. Intell.* **2015**, *29*, 508–514.
47. Tsai, A.C.-R.; Wu, C.-E.; Tsai, R.T.-H.; Hsu, J.Y.-j. Building a concept-level sentiment dictionary based on commonsense knowledge. *IEEE Intell. Syst.* **2013**, *28*, 22–30.
48. Cambria, E.; Livingstone, A.; Hussain, A. The Hourglass of Emotions. In *Cognitive Behavioural Systems*; Springer: Berlin/Heidelberg, Germany, 2012; pp. 144–157.
49. Susanto, Y.; Livingstone, A.G.; Ng, B.C.; Cambria, E. The Hourglass Model Revisited. *IEEE Intell. Syst.* **2020**, *35*, 96–102. [[CrossRef](#)]
50. Plutchik, R. The Nature of Emotions. *AmSci* **2001**, *89*, 344.
51. Glavaš, G.; Šnajder, J.; Bašić, B.D. Semi-Supervised Acquisition of Croatian Sentiment Lexicon. In *Text, Speech Dialogue, Proceedings of the 15th International Conference on Text, Speech and Dialogue, TSD 2012, Brno, Czech Republic, 3–7 September 2012*; Springer: Berlin/Heidelberg, Germany, 2012; pp. 166–173. [[CrossRef](#)]
52. Čoso B.; Guasch, M.; Ferré, P.; Hinojosa, J.A. Affective and concreteness norms for 3,022 Croatian words. *Q. J. Exp. Psychol.* **2019**, *72*, 2302–2312. [[CrossRef](#)]
53. Ilić, A.; Beliga, S. The Polarity of Croatian Online News Related to COVID-19: A First Insight. In Proceedings of the 32nd Central European Conference on Information and Intelligent Systems (CECIIS 2021), Varazdin, Croatia, 13–15 October 2021.

54. Agić, Ž.; Ljubešić, N.; Tadić, M. Towards Sentiment Analysis of Financial Texts in Croatian. *Bull. Mark.* **2010**, *143*, 69.
55. Babić, K.; Petrović, M.; Beliga, S.; Martinčić-Ipšić, S.; Matešić, M.; Meštrović, A. Characterisation of COVID-19-Related Tweets in the Croatian Language: Framework Based on the Cro-CoV-cseBERT Model. *Appl. Sci.* **2021**, *11*, 10442. [\[CrossRef\]](#)
56. Mršić, L.; Kopal, R.; Klepac, G. Analyzing Slavic Textual Sentiment Using Deep Convolutional Neural Networks. *Adv. Intell. Syst. Comput.* **2017**, *978*, 207–224.
57. Thakkar, G.; Mikelić Preradović, M.; Tadić, N. Multi-task Learning for Cross-Lingual Sentiment Analysis. *arXiv* **2022**, arXiv:2212.07160.
58. Rotim, L.; Šnajder, J. Comparison of Short-Text Sentiment Analysis Methods for Croatian. In Proceedings of the 6th Workshop on Balto-Slavic Natural Language Processing, Valencia, Spain, 4 April 2017; pp. 69–75. [\[CrossRef\]](#)
59. Goldberg, A.E. *Constructions: A Construction Grammar Approach to Argument Structure*; University of Chicago Press: Chicago, IL, USA, 1995.
60. Langacker, R. *Cognitive Grammar: A Basic Introduction*; Oxford University Press: New York, NY, USA, 2008.
61. Tomasello, M.; Brooks, P.J. Early syntactic development: A construction grammar approach. In *The Development of Language*; Psychology Press: Hove, UK, 1999; pp. 161–190.
62. Bergen, B.; Chang, N. Embodied Construction Grammar in Simulation-Based Language Understanding. In *Construction Grammars: Cognitive Grounding and Theoretical Extensions*; John Benjamins Publishing Company: Amsterdam, The Netherlands, 2005; Volume 3, pp. 147–190.
63. Sentiment Propagation. Available online: https://github.com/bperak/sentiment_propagation (accessed on 27 March 2025).
64. EnTenTen. Available online: https://app.sketchengine.eu/#dashboard?corpname=preloaded%2Fententen13_tt2_1 (accessed on 20 March 2025).
65. Csardi, G.; Nepusz, T. The igraph Software Package for Complex Network Research. *InterJournal* **2006**, *Complex Systems*, 1695. Available online: <https://igraph.org> (accessed on 20 March 2025).
66. hrWac22. Available online: https://app.sketchengine.eu/#dashboard?corpname=preloaded%2Fhrwac22_ws (accessed on 20 March 2025).
67. De Schryver, G.M. Generative AI and Lexicography: The Current State of the Art Using ChatGPT. *Int. J. Lexicogr.* **2023**, *36*, 355–387.
68. Vaswani, A.; Shazeer, N.; Parmar, N.; Uszkoreit, J.; Jones, L.; Gomez, A.N.; Kaiser, L.; Polosukhin, I. Attention Is All You Need. *arXiv* **2017**, arXiv:1706.03762.
69. Höder, S. Grammar is Community-Specific: Background and Basic Concepts of Diasystematic Construction Grammar. In *Constructions in Contact: Constructional Perspectives on Contact Phenomena in Germanic Languages*; Boas, H.C., Höder, S., Eds.; John Benjamins Publishing Company: Amsterdam, The Netherlands, 2018; pp. 37–70. [\[CrossRef\]](#)
70. Liu, Z.; Qiao, A.; Neiswanger, W.; Wang, H.; Tan, B.; Tao, T.; Li, J.; Wang, Y.; Sun, S.; Pangarkar, O.; et al. LLM360: Towards Fully Transparent Open-Source LLMs. *arXiv* **2023**, arXiv:2312.06550.

Disclaimer/Publisher’s Note: The statements, opinions and data contained in all publications are solely those of the individual author(s) and contributor(s) and not of MDPI and/or the editor(s). MDPI and/or the editor(s) disclaim responsibility for any injury to people or property resulting from any ideas, methods, instructions or products referred to in the content.