

Automating Adjectival Microstructures in Monolingual Dictionaries: A New Method Combining Embeddings and LLMs

Enikő Héja¹, László Simon¹, Veronika Lipp¹

¹ ELTE Research Centre for Linguistics, Budapest, 1068 Benczúr u. 33. Hungary
E-mail: heja.eniko@nytud.hu, simon.laszlo@nytud.hu, lipp.veronika@nytud.hu

Abstract

Recent findings indicate that current large language models (LLMs) face difficulties in generating clear-cut, well-motivated definitions in a consistent way. This shortcoming is the consequence of their reliance on opaque data sources and their inherently unstable, non-deterministic outputs. In response, this research aims to develop an LLM-based methodology for producing adjectival microstructures in monolingual dictionaries in a way that is both more consistent and aligned with lexicographic standards. Building on the hypothesis that prompts enriched with contextual information can enhance definition quality, the study employs a graph-based, interpretable, and unsupervised method starting out from static adjectival embeddings. The approach has previously demonstrated the ability to formalize traditional lexical semantic relations, detect adjectival senses from corpus data, and identify the most salient nominal contexts for each sense. The ultimate goal is to integrate these results into practical lexicographic workflows and assess how LLMs, when properly guided, can support dictionary compilation.

Keywords: unsupervised sense detection; graphs; adjectival microstructure; LLMs

1. Introduction

In recent decades, advances in semantic theory—especially cognitive and frame semantics—have significantly influenced lexicography, shaping how word meanings are described (Ostermann, 2015). This has led to the development of a wide array of lexical resources, from dictionaries to word embeddings and semantic networks. These tools have not only enriched semantic research but also expanded lexicography’s role in applied fields such as information retrieval and the development of Large Language Models (LLMs), highlighting the growing synergy between lexicography and semantics (Štrkalj Despot et al., 2024).

However, recent findings have shown that current LLMs struggle to produce clear-cut, well-motivated definitions, partly, because they are based on “*unknown data sources, with non-deterministic (and very likely soon-to-be-personalized) responses, very limited stability and reproducibility*” (Jakubíček and Rundell, 2023).

Thus, the primary goal of the research is to introduce an LLM-based methodology that is able to automatically prepare the adjectival microstructures in monolingual dictionaries in a more deterministic and consistent way. The present research is based on the hypothesis that complementing prompts with abundant contextual information, will help us to generate lexicographically sound adjectival entries.

Our methodology starts from the unsupervised, data-driven exploration of adjectival lexical semantics. For that purpose, an interpretable graph-based method is applied (Héja et al., 2022a; 2023; 2024a; 2024b) that utilizes adjectival static embeddings. Based on our proof-of-concept experiments, the proposed method is able (1) to grasp traditional lexical semantic relations in a more formalized way, (2) to detect adjectival senses on the basis of corpus data, (3) to identify the most relevant nominal contexts specific to each sense. (4) We intend to put the results into lexicographic practice, and to investigate to what extent the method is able to support the compilation of a monolingual dictionary, when complemented with LLMs.

For that purpose, we selected a set of adjectives—*fekete* ‘black’, *magas* ‘high’, *mély* ‘deep’, *szabad* ‘free’, *könnyű* ‘easy’, *sötét* ‘dark’, and *nagy* ‘great’—which share linguistically and lexicographically relevant features. These adjectives are all polysemous, with multiple subsenses spanning literal, abstract, and metaphorical domains. Their meanings are highly context-dependent, and they frequently occur with specific noun types that help signal particular senses (e.g., *fekete ruha* ‘black clothes’ vs. *fekete piac* ‘black market’ vs. *fekete történelem* ‘dark history’). This context sensitivity makes them well-suited for distributional modelling. Due to their semantic complexity, they pose challenges for dictionary compilation—particularly in sense disambiguation and example selection—thus providing an ideal testing ground for the approaches.

We believe that the detected subsenses along with their contextual clues in the prompts may eliminate the above mentioned deficiencies. For instance, whereas prompt1 produced a somewhat arbitrary and inconsistent answer, prompt2, complemented with the distributional knowledge from our graph-based method, made ChatGPT create a lexicographically sound entry.

Prompt1: “Please write definitions in Hungarian to the subsenses of the adjective ‘Ferences [Franciscan]’.”

Prompt2: “Please write definitions in Hungarian to the subsenses of the adjective ‘Ferences’ if you know that typically {‘rendház’, ‘monostor’, ‘kolostor’}¹; {‘pap’, ‘szerzetes’, ‘apát’}² and {‘egyetem’, ‘gimnázium’, ‘iskola’}³ can be ‘Ferences’. Give corresponding example sentences in Hungarian as well.”

¹ ‘friary’, ‘monastery’, ‘convent’, respectively.

² ‘priest’, ‘monk’, ‘abbot’, respectively.

³ ‘university’, ‘secondary school’, ‘school’, respectively.

2. Objectives of the Research

The primary goal of this research is to explore adjectival lexical semantics in an unsupervised, data-driven manner. To this end, we plan to enhance an existing proof-of-concept methodology, ensuring it encompasses all sufficiently frequent adjectives in our corpus while improving the robustness of lexical semantic relation detection. This graph-based method, utilizing adjectival static embeddings, not only identifies adjectives with specific semantic relationships but also links them to particular contexts. Although our focus is on adjectives, the proposed methodology is not inherently restricted to this part of speech; it can be extended to other classes of content words, such as nouns and verbs, for detecting their lexical-semantic structures as well. Additionally, we aim to evaluate our method’s effectiveness by comparing its output to an existing dictionary and also to assess its potential contribution to the compilation process of monolingual dictionaries.

Our research is motivated by the fact that, despite some existing studies on Hungarian adjectives (e.g., Kiefer, 2000; 2008), there is no comprehensive description of adjectives in the scientific literature. Moreover, previous studies (cf. Apresjan, 1974; Haber and Poesio, 2024; Ježek, 2016; Pustejovsky, 1995 etc.) reveal significant inconsistencies in defining key lexical-semantic relations such as homonymy and various types of polysemy (regular, irregular, systematic), etc. These inconsistencies make it challenging to apply such definitions in practice. Therefore, the methodology we plan to develop should clarify these fundamental concepts and should be able to cut through the terminological confusion, helping us also to develop a clear understanding of these basic notions.

This conceptual ambiguity also impacts lexicographic practice, causing inconsistencies in both the overall structure and in the detailed entries of dictionaries (cf. Héja et al., 2024a; Svensén, 2009). To address this, we will test our approach through comparing it to the Concise Dictionary of Hungarian (2003) and examine its potential use in compiling the Comprehensive Dictionary of Hungarian (2006).

In summary, we aim to develop a solid and robust methodology that is able (1) to grasp traditional lexical semantic relations in a more formalized way, (2.a.) to detect adjectival senses on the basis of corpus data, (2.b.) to identify the contexts specific to each sense. We also plan to (4) validate our results on the basis of a monolingual dictionary. (5) Finally, we intend to put the results into practice, and apply the results to support the compilation of the Comprehensive Dictionary of Hungarian utilizing the most recent advances of language technology, such as large language models (LLMs henceforward).

3. Methodology

The proposed algorithm will be an improved version of our proof-of-concept method (Héja and Ligeti-Nagy, 2022a; 2022b), a graph-based distributional approach originally inspired by the work of Ah-Pine and Jacquet (2009). This methodology has been shown to effectively extract monolingual adjectival lexical semantic relations, such as synonymy and systematic polysemy, for more than 10,000 adjectives directly from corpus data (Héja et al., 2023; Héja et al., 2024a; Héja et al., 2024b).

3.1 Lexical semantic relations as local graph structures

In essence, it was found that, if a suitable representation of adjectives is converted into a suitable graph G , certain local graph structures of G naturally reveal lexical-semantic relationships. For instance, a sense—conceived of as synonyms forming a distinct synonym set (in the sense of WordNet, Miller, 1995)—can be identified as complete adjectival subgraphs, or cliques (Fig. 1). Likewise, polysemous adjectives (Fig. 2) can be recognized as those belonging to multiple synonym sets, meaning they appear in multiple cliques. The following automatically detected graph structures illustrate these patterns.

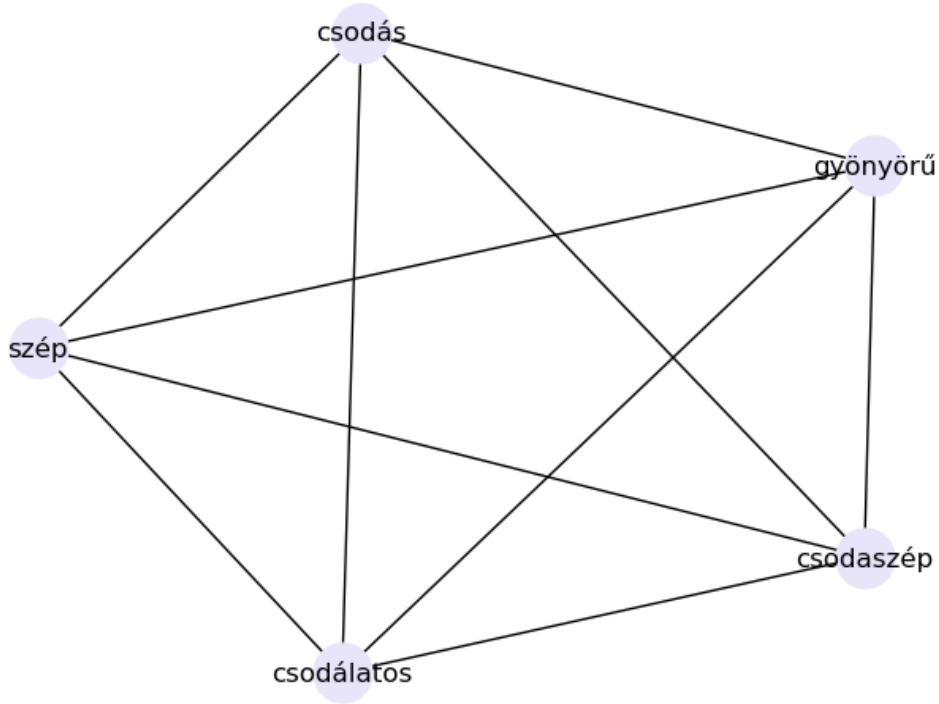


Fig. 1: Synonym set of ‘szép’ (*nice*)⁴

⁴ szép 'beautiful'; csodás 'wonderful'; gyönyörű 'gorgeous'; csodaszép 'stunning'; csodálatos 'magnificent'

One particularly intriguing feature of the automatically generated adjectival graph is its ability to group all adjectives in the corpus into semantically coherent fields. This is achieved by identifying the connected components within the graph. For example, the first graph (Fig. 3) represents adjectives related to monastic orders, while the second (Fig. 4) consists of adjectives derived from country names.

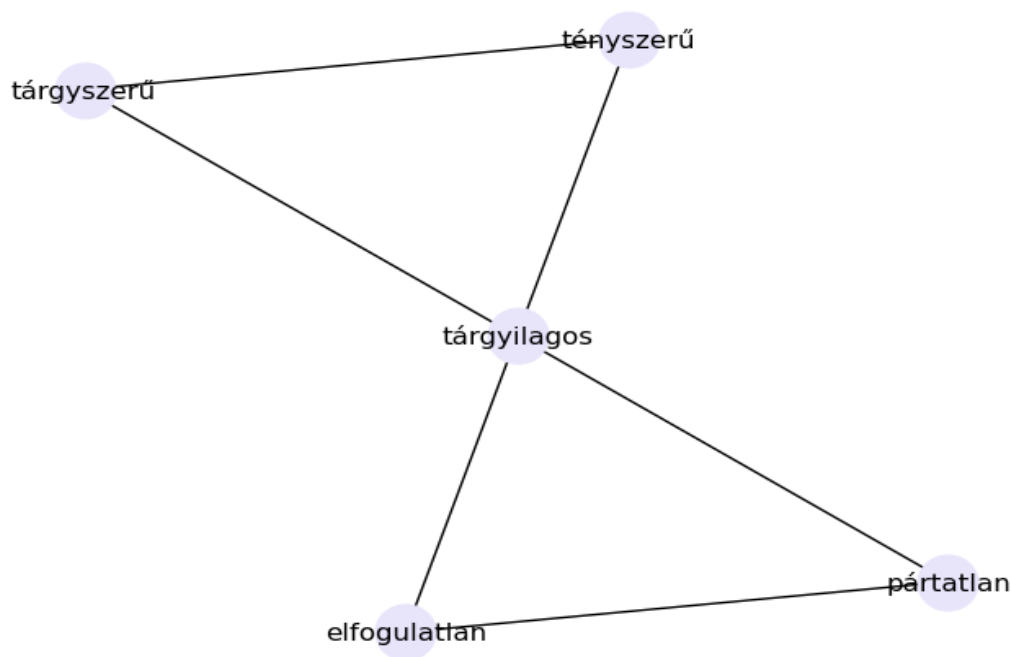
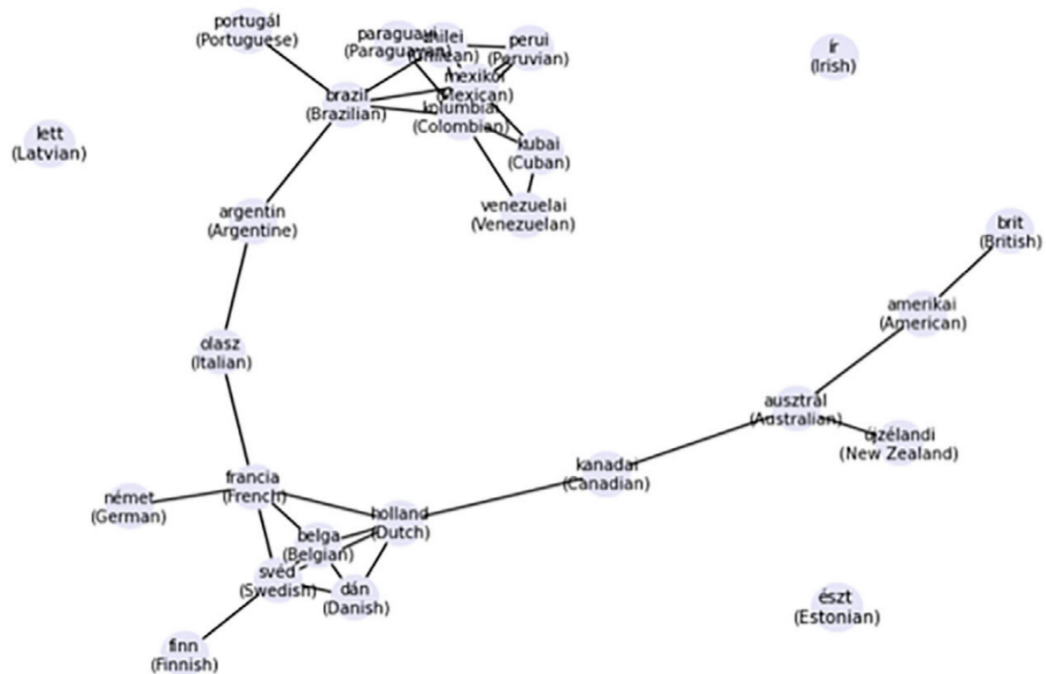
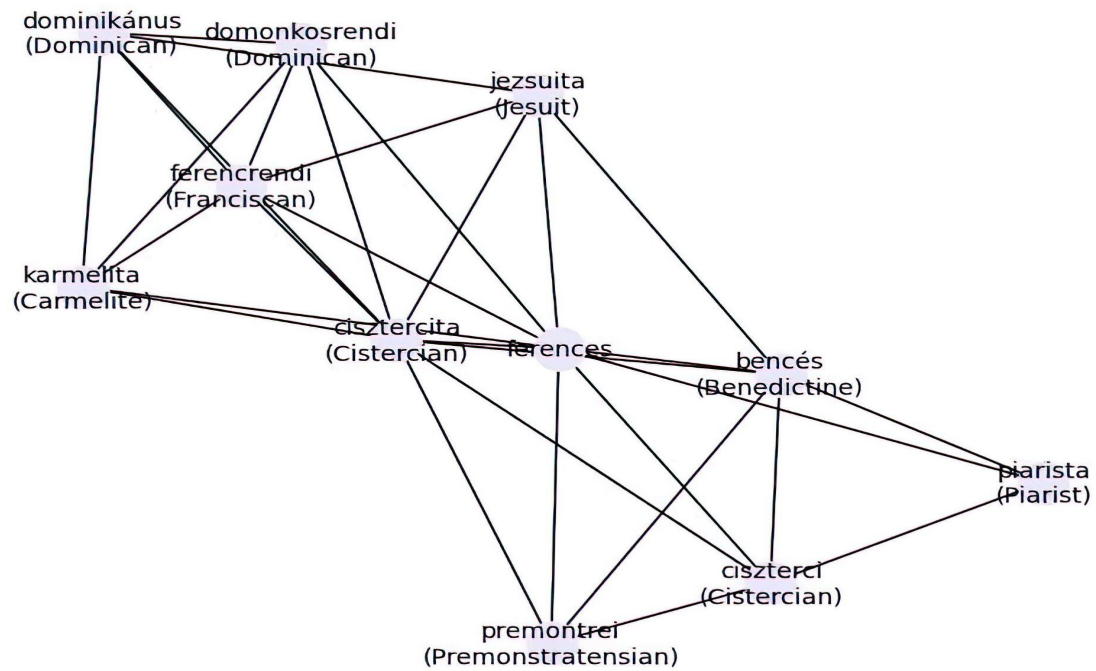


Fig. 2: Polysemous senses of ‘tárgyilagos’ (*objective*)⁵

⁵ tárgyszerű 'objective'; tényyszerű 'factual'; tárgyilagos 'impartial'; pártatlan 'neutral'; elfogulatlan 'unbiased'



Our findings also suggest that the proposed methodology can detect homonyms, provided we define homonymy as the accidental collision of two random senses (as opposed to less operationalizable definitions, such as those based on etymology or based on the property of sense unrelatedness). In this framework—due their unique distributions—the homonymous word forms appear as *isolated nodes* in the adjectival graph, positioned far from all other nodes. Fig. 4. clearly shows that homonymous country-name adjectives—such as *lett* (‘Latvian’ vs. ‘became’), *ír* (‘Irish’ vs. ‘write’), and *észt* (‘Estonian’ vs. ‘intelligence’ in accusative form)—emerge as isolated nodes within the COUNTRY subgraph of G .

3.2 Construction of the adjectival graph

Here, we briefly summarize the technical workflow (for a more detailed explanation, see Héja and Ligeti-Nagy, 2022b and Zinoviev, 2018). The adjectival graph is built using static word embeddings, which capture the distributional properties of adjectives in corpus data (e.g., Mikolov et al., 2013; Rehurek & Sojka, 2011). First, an adjacency matrix (A) of size $N \times N$ is generated, where N represents the size of the adjectival vocabulary. Each matrix cell $\{i, j\}$ is computed using a similarity metric—such as cosine similarity—between the vector representations of adjectives Adj_i and Adj_j . Next, a weighted, undirected graph (W) is constructed based on matrix A , where N nodes (adjectives) are fully interconnected. The edges in W are assigned weights reflecting the degree of semantic similarity between adjective pairs. To create a binary (unweighted) graph (G), we apply a thresholding process to W , using a suitable K cut-off parameter.

Finally, we extract key graph structures that likely correspond to specific lexical-semantic relations. These include: cliques for synonym sets representing senses, shared nodes in multiple cliques for polysemous adjectives, isolated nodes for homonyms and connected components for semantic fields. This extracted information provides insights into the organization of adjectival lexical semantics on the basis of corpus data.

3.3 Retrieving the relevant contexts

In the final step of the algorithm, relevant nominal contexts are extracted to determine the most salient contexts for each specific adjectival sense. While the extraction process varies depending on the lexical-semantic relation (e.g., polysemy, homonymy, systematic polysemy, etc.), it consistently relies on unsupervised clustering methods to identify the most relevant contexts without manual intervention.

From a lexicographic perspective, the most valuable outcome of this process is the identification of a data structure, which we call *adjectival meaning structures*. These structures, derived from corpus data, capture systematic (or regular) polysemies in adjectives—at least if we accept Apresjan’s definition of regular polysemy (1974), and therefore, are particularly useful constructions from the perspective of lexicography. For example, Fig. 5 illustrates a segment of the semantic meaning structure of monastic orders. It states that *bencés* ‘Benedictine’, *ciszterci*, ‘Cistercian’ etc. are all monastic

orders; monastic orders tend to own buildings (such as monasteries or convents), tend to have positions in their organisations (such as monks, abbots and priests) and tend to maintain educational organisations (such as universities or high schools).

The numbers in the meaning structure indicate that the properties identified by the columns represent tendencies rather than absolute rules, as not all monastic orders exhibit all of these characteristics. It is important to note that adjectival meaning structures are derived not from individual adjectives, but from semantic domains, enriched with their most relevant contexts constituting the columns of the meaning structures. These contexts are identified fully automatically by clustering and co-occurrence analysis (arithmetic mean), ensuring a data-driven approach to understanding adjectival meaning patterns.

bencés ‘Benedictine’	182	98.14	87.6
ciszterci ‘Cistercian’	49.29	26.43	331
cisztercita ‘Cistercian (fem)’	20.83	8.5	
dominikánus ‘Dominican’	32	11.83	
domonkosrendi ‘Dominican’	1.75		
ferences ‘Franciscan’	233.67	192.71	76.75
ferencrendi ‘Franciscan’	4.67	6.75	
jezsuita ‘Jesuit’	49.8	172.86	88.6
karmelita ‘Carmelite’	48.33	12.4	
piarista ‘Piarist’	49.67	94.6	193
premontrei ‘Premonstratensian’	36.88	20.17	73.5
	rendház ‘monastery’	szerzetes ‘monk’	iskola ‘school’
	monostor ‘monastery’	apát ‘abbot’	egyetem ‘university’
	kolostor ‘convent’	pap ‘priest’	gimnázium ‘high school’

Fig. 5: Part of the meaning structure of the lexical set ‘monastic orders’

4. Expected Results

Our proof-of-concept experiments, conducted in collaboration with expert lexicographers, have demonstrated that ChatGPT complemented with structured data, which was extracted from large corpora, can be used for lexicographic purposes, resulting in consistent micro- and macrostructure.

4.1 Structuring dictionary entries (microstructure level)

By retrieving adjectival conceptual structures, our method helps to create microstructure templates, ensuring consistency across adjectival dictionary entries. For instance, all monastic order adjectives (*bencés* ‘Benedictine’, *ciszterci* ‘Cistercian’, etc.) share the same set of semantic components in their microstructures, maintaining uniformity in lexical descriptions.

4.2 Categorizing adjectival vocabulary (macrostructure level)

By effectively grouping adjectives into well-defined semantic categories, our approach provides insights into productive adjectival derivation and supports the development of more systematic principles for dictionary editing. For example, adjectives like *karnyújtásnyi* (‘within arm’s reach’), *kéznyújtásnyi* (‘within hand’s reach’), *kőhajításnyi* (‘within a stone’s throw’), *nyíllövésnyi* (‘within an arrow’s shot’), and *puskalövésnyi* (‘within a gunshot’) form a cohesive graph component. Their strong structural similarity raises lexicographic questions—should they all be listed as individual headwords, or should they be treated as a derivational pattern and omitted altogether? However, if one of them is included, consistency demands that the entire group be included, yet this is not the case in the Concise Dictionary of Hungarian.

4.3 Automating extraction of good dictionary examples (GDEXs)

Using the automatically distinguished sense candidates and their associated nominal contexts, we can extract corpus-based dictionary examples (GDEX) that effectively illustrate word usage (Kilgariff et al., 2008). This allows for the automatic generation of high-quality example sentences for lexicographic entries.

4.4 Enhancing definitions with LLMs

As part of our research, we plan to leverage LLMs to complement the extracted adjectival meaning structures with automatically generated definitions, using fine-tuned SOTA models. We believe that our predefined meaning structures, derived from high-quality data, can help address the limitations mentioned in the Introduction (cf. Jakubíček and Rundell, 2023). By refining existing LLMs through prompt engineering or fine-tuning, we will be able to generate accurate and lexicographically sound adjectival entries.

In the next section, we provide a detailed evaluation of how the proposed method performs in practical lexicographic tasks, with a focus on 4.1, 4.2 and 4.3. This includes an assessment of how effectively ChatGPT, when guided by our meaning structures, can produce dictionary-quality entries for a representative sample of adjectives.

5. Evaluation

Output 1 involved the manual revision of five newly created meaning structures presented in Excel tables. These tables were reviewed by lexicographers and included only columns relevant from a lexicographic perspective. As a result, unnecessary information was removed, and the structure was better aligned with the requirements of dictionary-making.

Output 2 was based on these refined tables and involved the generation of full dictionary entries using ChatGPT. The first prompt was in English:

Prompt 1:

“Please write definitions in Hungarian to the subsenses of the adjective 'mély'⁶ if you know that typically {1: ['mondanivaló'], 2: ['elemzés', 'jelentés'], 3: ['összefüggés'], 4: ['szín'], 5: ['magánhangzó'], 6: ['víz', 'sár', 'talaj', 'homok'], 7: ['kátyú', 'nyomvályú', 'úthiba'], 8: ['hasadék', 'lyuk', 'ránc', 'repedés', 'vágás', 'üreg', 'bevágás', 'barázda'], 9: ['tenger', 'völgy', 'meder', 'akna', 'szurdok', 'tó', 'vízmosás', 'öböl', 'barlang', 'medence', 'kút', 'kanyon', 'gödör', 'munkagödör', 'kráter', 'árok'], 10: ['pince', 'út', 'csatorna', 'folyó', 'verem', 'gyökér'], 11: ['erdő', 'vágás', 'légzés', 'seb', 'késszúrás'], 12: ['tál', 'tányér'], 13: ['zseb', 'bugyor'], 14: ['fotel'], 15: ['pont'], 16: ['hó'], 17: ['probléma', 'megosztottság', 'ellentét', 'konfliktus'], 18: ['egyetértés'], 19: ['szegénység', 'nyomor'], 20: ['úr', 'szakadék'], 21: ['átalakítás', 'változás'], 22: ['válság', 'hullámvölgy', 'recesszió'], 23: ['barátság', 'kapcsolat', 'kötődés'], 24: ['érdeklődés'], 25: ['felelősség'], 26: ['meditáció', 'légzés', 'alvás', 'altatás'], 27: ['torok', 'seb', 'fájdalom'], 28: ['dekoltázs'], 29: ['öröm', 'élmény', 'csalódás', 'érzés'], 30: ['dolog', 'gondolat'], 31: ['benyomás'], 32: ['aggodalom', 'vágy', 'megbánás', 'depresszió', 'bánat', 'érzelem', 'csalódottság', 'ellenszenv', 'letargia', 'szerelem', 'résztét', 'szomorúság', 'megvetés', 'megdöbbenés', 'gyász', 'bizalmatlanság', 'vonzalom', 'felháborodás', 'sajnálkozás', 'szenvédély', 'megrendülés', 'sajnálat', 'csodálat', 'nyugtalanság', 'rokonszenv', 'gyűlölet', 'magány'], 33: ['tudás', 'élettapasztalat', 'ismeret'], 34: ['vallásosság', 'meggyőződés', 'elkötelezettség'], 35: ['tisztelet', 'hódolat'], 36: ['átézés', 'szeretet', 'együttérés', 'bölcsség', 'alázat', 'megértés', 'átélés', 'értelem', 'emberség', 'hit', 'empátia', 'alázatosság', 'igazság'], 37: ['hangzás', 'zene', 'basszus', 'hang'], 38: ['torokhang', 'férfihang'], 39: ['sötétség', 'homály', 'csend', 'árnyék', 'csönd'], 40: ['álom'], 41: ['beszélgetés', 'áhítat'], 42: ['lélegzetvétel', 'sóhaj', 'lélegzet', 'sóhajtás'], 43: ['hallgatás', 'meghajlás'], 44: ['tekintet'], 45: ['hála']}]⁷ can be 'mély'. Each number in the list along with the nouns in square brackets refers to a specific sub-sense. Give corresponding example sentences in Hungarian as well. Please be precise, confine yourself to the nouns listed, both when writing the definitions and

⁶ deep

⁷ 1: ['message'] 2: ['analysis', 'report'] 3: ['connection'] 4: ['color'] 5: ['vowel'] 6: ['water', 'mud', 'soil', 'sand'] 7: ['pothole', 'rut', 'road defect'] 8: ['crack', 'hole', 'wrinkle', 'fissure', 'cut', 'cavity', 'notch', 'furrow'] 9: ['sea', 'valley', 'riverbed', 'shaft', 'gorge', 'lake', 'gully', 'bay', 'cave', 'basin', 'well', 'canyon', 'pit', 'work pit', 'crater', 'ditch'] 10: ['cellar', 'road', 'sewer', 'river', 'pit', 'root'] 11: ['forest', 'clearing', 'breath', 'wound', 'stab wound'] 12: ['bowl', 'plate'] 13: ['pocket', 'pouch'] 14: ['armchair'] 15: ['point'] 16: ['snow'] 17: ['problem', 'division', 'opposition', 'conflict'] 18: ['agreement'] 19: ['poverty', 'misery'] 20: ['void', 'abyss'] 21: ['transformation', 'change'] 22: ['crisis', 'downturn', 'recession'] 23: ['friendship', 'relationship', 'attachment'] 24: ['interest'] 25: ['responsibility'] 26: ['meditation', 'breathing', 'sleep', 'anesthesia'] 27: ['throat', 'wound', 'pain'] 28: ['cleavage'] 29: ['joy', 'experience', 'disappointment', 'feeling'] 30: ['thing', 'thought'] 31: ['impression'] 32: ['worry', 'desire', 'regret', 'depression', 'sorrow', 'emotion', 'disappointment', 'antipathy', 'lethargy', 'love', 'compassion', 'sadness', 'contempt', 'astonishment', 'grief', 'distrust', 'attraction', 'indignation', 'lament', 'passion', 'distress', 'pity', 'admiration', 'uneasiness', 'sympathy', 'hatred', 'loneliness'] 33: ['knowledge', 'life experience', 'awareness'] 34: ['religiosity', 'conviction', 'commitment'] 35: ['respect', 'homage'] 36: ['sympathy', 'love', 'compassion', 'wisdom', 'humility', 'understanding', 'empathy', 'reason', 'humanity', 'faith', 'empathy', 'humbleness', 'truth'] 37: ['sound', 'music', 'bass', 'voice'] 38: ['throaty voice', 'male voice'] 39: ['darkness', 'dimness', 'silence', 'shadow', 'quiet'] 40: ['dream'] 41: ['conversation', 'devotion'] 42: ['breath', 'sigh', 'breathing', 'sighing'] 43: ['silence', 'bowing'] 44: ['gaze'] 45: ['gratitude']

also when you create the relevant example sentences. Do not use the headword ('mély') in the definition. Be also comprehensive and cover all the listed sub-senses with definitions."

In the next phase, the prompt was rewritten in Hungarian and enriched with more detailed guidance, designed to improve the quality, clarity, and consistency of the dictionary entries. The revised prompt asked the model to write definitions in Hungarian for each subsense of *mély*, based on noun groups associated with that meaning. Each number in curly brackets represented a distinct subsense, and the nouns listed in square brackets served as contextual anchors for the interpretation. The model was instructed to provide a precise Hungarian definition and a corresponding example sentence that included one of the associated nouns. Crucially, the adjective *mély* was explicitly forbidden in the definition but was required in every example sentence to ensure that the definition would not simply repeat or circularly reference the headword. To support proper interpretation, the prompt included examples of acceptable and unacceptable formulations. For instance, a noun-based (and thus incorrect) interpretation was: "in contrast with white: dark-colored piece, rook". The correct adjectival formulation would be: "in contrast with white: black or dark brown in color." Similarly, "not adequately illuminated: dark forest, yard" was deemed incorrect, while "not adequately illuminated" was accepted as a valid adjectival interpretation.

Prompt 2:

"Kérlek, írd meg magyar nyelvű definíciókat a könnyű melléknév magyar jelentéseit illusztrálva. A kapcsos zárójelek közötti listában minden szám egy-egy külön jelentésárnyalatot jelöl, és a hozzá tartozó szögletes zárójelben felsorolt főnevek alapján kell meghatározni, hogy a könnyű melléknévnek mi a jelentése akkor, amikor felsorolt főnevek előtt áll. Adj pontos magyar nyelvű definíciókat minden egyes jelentéshez, és mindig adj meg egy jó magyar nyelvű példamondatot is. Olyan példamondatot, amelyben szerepel az adott sorszámhoz tartozó főnév vagy főnevek valamelyike. A definíciókban a könnyű szó nem szerepelhet. A példamondatokban a könnyű szónak mindig szerepelnie kell. Példák a jó és rossz értelmezésekre: Főnévi, rossz értelmezés, a sötét szó használatával: 'a világossal ellentétben: sötét színű figura, bástya, futó' Melléknévi, azaz jó értelmezés: 'a világossal ellentétben: fekete v. sötétbarna színű' Főnévi, rossz értelmezés, a sötét szó használatával: 'kellően meg nem világított, sötét erdő, udvar' Melléknévi, azaz jó értelmezés: 'kellően meg nem világított' [...]." [Please write Hungarian-language definitions to illustrate the meanings of the adjective könnyű. Each number in the list enclosed in curly brackets represents a distinct sub-sense, and the nouns listed in square brackets specify which sense of könnyű should be defined. Your task is to determine what könnyű means when it modifies each of the listed nouns. For every sub-sense, provide a precise definition in Hungarian, followed by a well-formed Hungarian example sentence that includes the adjective könnyű and one of the corresponding nouns. In your definitions, the word könnyű must not appear. The definitions should be purely adjectival in nature (i.e., capable of being substituted for the adjective in the sentence), and should not be noun-like explanations. The example sentence must contain the word könnyű, and the noun

associated with the relevant subsense must appear in the sentence as well. To clarify what kind of definitions are appropriate, here are some examples of incorrect and correct interpretations using the adjective *sötét*. An incorrect (noun-type) definition would be: “a világossal ellentétben: sötét színű figura, bástya, futó.” A correct (adjectival) definition would be: “a világossal ellentétben: fekete vagy sötétbarna színű.” Another incorrect (noun-type) definition: “kellően meg nem világított, sötét erdő, udvar.” The correct adjectival interpretation would be: “kellően meg nem világított.”]

This more comprehensive prompt helped the model produce clearer, more lexicographically sound definitions. The ChatGPT-generated entries were subjected to a separate evaluation to assess their quality, structure, and fidelity to the underlying data.

5.1 Evaluation of *sötét* (adj.) – dark, dim, lacking light or brightness, figuratively obscure or grim, depending on context

The meaning structure of *sötét* (‘dark’) outlined in the initial table corresponds closely to the structure found in the Concise Dictionary of Hungarian, with one notable exception: the table separates certain groups of meanings that are subsumed under a single definition in the dictionary. For example, “dark eyes,” “dark skin,” “dark trousers,” and “dark colour” are all grouped under definition 2. in the dictionary: “a shade tending towards or resembling black.”

After manually revising the 61 noun groups listed in the table—which had been used as input for a ChatGPT prompt—we identified 43 subsenses. Out of these, 11 were hallucinated, 3 represented novel senses not documented in the Concise Dictionary of Hungarian. (eg. ‘<A period> filled with hardship, suffering, or danger.’; ‘<A galaxy> that emits little or no visible light.’; ‘<A type of matter> that cannot be directly observed with astronomical instruments.’) and 27 could be clearly aligned with dictionary definitions.

For Prompt 2 ChatGPT successfully generated adjectival senses and, in 5 cases, the definitions were more precise than in previous attempts. Out of the 43 subsenses, three errors occurred in which *sötét* was used in the definition despite the explicit instruction, and in one case the example sentence did not correspond to the given meaning.

5.2 Evaluation of *könnyű* (adj.) – light, easy, effortless, low in intensity, depending on context

The spreadsheet identified 61 distinct noun groups associated with the adjective *könnyű* (‘light’ or ‘easy’). When these were compared to the meaning structure in the Concise Dictionary of Hungarian, 32 groups could be directly matched to dictionary entries. Notably, the table also introduced 4 additional submeanings that are not present in the printed source.

For the Prompt 1 the ChatGPT produced a remarkably accurate set of definitions, distinguishing 32 subsenses in total. Of these, only one was a hallucination, while the remaining definitions aligned well with the dictionary structure. For the 4 novel meanings, ChatGPT offered precise and contextually appropriate definitions — for instance, *könnyű drog* was defined as “a low-potency, mild psychoactive substance.” In keeping with the instructions, the model avoided using the headword *könnyű* in its definitions. While it did occasionally use the inflected form *könnyen* (‘easily’), this was not explicitly forbidden by the prompt and therefore does not count as an error.

At the next stage, the results for Prompt 2 resembled those of the earlier prompt; however, in this case the model failed to identify two meanings it had previously recognized without difficulty.

5.3 Evaluation of *magas* (adj.) – tall, high, elevated, significant, depending on context.

The table identified 79 semantic groups, of which 58 were retained after manual analysis. Each of these 58 groups could be aligned with the meaning structure found in the Concise Dictionary of Hungarian, except for one missing sense that did not emerge from the vector-based table: 3. ‘Of high frequency <sound>. Music: ~ c: the c note near the upper limit of a soprano or tenor singer’s vocal range.’

In response to the English-language prompt, ChatGPT produced accurate adjectival definitions; in accordance with the instructions, the adjective *magas* was not repeated in the definitions, the example sentences were all appropriate, and in one case, even the superlative form of the adjective was used:

Sense 36. [‘pozíció’, ‘kulcs’, ‘funkció’]⁸

Fontosságában vagy hierarchiában kiemelkedő. [Outstanding in importance or hierarchical level.] Example: Ő tölti be a vállalat legmagasabb pozícióját. ‘She holds the highest position in the company.’)

Subsequently, we examined the 58 dictionary entry structures in response to the Hungarian-language prompt with more detailed instructions. The first 4 entries matched exactly with the ones produced from the English prompt. However, starting with the fifth, ChatGPT began to slightly modify the example sentences. For instance, for sense 14, the definition was identical in both cases (‘Összegszerűségében jelentős.’ – ‘Significant in amount.’), but the example sentences differed slightly:

- From the first prompt: *A skandináv országokban magas a nyugdíj és a kereset.* [In the Scandinavian countries, pensions and earnings are high.]

⁸ ‘position’, ‘key’, ‘function’, respectively.

- From the second prompt: *Skandináviában különösen magas az adóteher.* [In Scandinavia, the tax burden is particularly high.]

Similarly, for sense 28, the definition remained the same in both versions (‘Árban jelentős, megterhelő.’ – ‘Significant or burdensome in price.’), but the example sentences varied:

- Example 1: *A gazdák szerint idén különösen magas a takarmányár.* [According to the farmers, the price of fodder is particularly high this year.]
- Example 2: *A gazdák panaszkodnak a magas takarmányárra.* [The farmers are complaining about the high price of fodder.]

5.4 Evaluation of *mély* (adj.) – deep, low, intense, profound, depending on context

The table identified 65 semantic groups, of which 45 were retained after manual analysis. Each of these 45 groups could be aligned with the meaning structure found in the Concise Dictionary of Hungarian, except for one missing sense that did not emerge from the vector-based table: 6. sense: ‘At a low level in terms of material or moral values.’

For prompt 1 in English, the resulting dictionary entries were excellent, such as:

‘*mély*’ 1. *Tartalmában komoly, gondolatébresztő.* [Serious in content, thought-provoking.]

Example: *A beszéd mély mondanivalót hordozott.* [The speech carried a deep message.]

or a more specific meaning:

‘*mély*’ 5. *A szájüreg alsó részében képzett, ajakréses vagy hátul képzett.* [Formed in the lower part of the oral cavity, with lip rounding or produced at the back.]

Example: *Az „a” és „o” a magyar nyelv mély magánhangzói közé tartoznak.* [The sounds “a” and “o” are among the deep vowels of the Hungarian language.]

This time, all definitions were adjectival interpretations, and all example sentences included the headword.

For Prompt 2, the model provided a more precise definition in some cases. For example, in response to Prompt 1, it gave the following:

28. ['dekoltázs']

Downward-sloping, revealing a lot.

Example: The dress was designed with a deep neckline.

Whereas for Prompt 2, it returned:

28. ['dekoltázs']

Meaning: Long, deeply cut (clothing element).

Example sentence: The actress wore an elegant dress with a deep neckline.

5.5 Evaluation of *nagy* (adj.) – big, large, great, significant, depending on context

Of the 97 semantic groups outlined in the table, 71 correspond to sense 5 in the Concise Dictionary of Hungarian (“Of considerable extent or degree. e.g. serious trouble, heavy storm. | Important, significant.”). One sense did not appear in the table: sense 7 – “Loud-sounding but lacking substance.” (e.g. grandiose words), which is not an error, as this expression functions as a fixed collocation rather than a compositional adjective–noun pairing.

As for the entries generated in response to the first prompt, their content was characterized by the fact that, although most noun groups fell under the same dictionary sense (i.e. the 5th sense), ChatGPT attempted to provide a separate definition for each noun group. This occurred because the model treated each nominal cluster as indicative of a distinct subsense, even when, from a lexicographic point of view, they could have been grouped under a single overarching meaning. For example:

(a) Typical nouns: nyomorúság, pusztítás, tűzvész, támadás, csapás, megrázkódtatás, tragédia, katasztrófa, földrengés, kolerajárvány, árvíz, szenvedés, roham, megpróbáltatás, robbanás, éhínség, trauma⁹
Definition: An event that is destructive in its impact and entails serious consequences.

(b) Typical nouns: titok, kincs, téma, talány, kérdőjel, dilemma, rejtély, kérdés, tanulság¹⁰
Definition: A phenomenon or concept that is particularly profound or thought-provoking in its content or significance.

When constructing the definitions, the model was only able to provide adjectival interpretations in two cases (cf. “Intense, dynamic in movement or emotional state”; “Notable in emotional impact or significance”); in the rest, it produced nominal definitions, which is an error.

In response to the Hungarian-language prompt, which included more detailed instructions, ChatGPT managed to provide adjectival definitions in all but one case, and all example sentences were appropriate. However, a consistent pattern across all adjectives examined was that the model rarely used inflected forms of the adjective in

⁹ 'misery', 'destruction', 'conflagration', 'attack', 'blow', 'shock', 'tragedy', 'catastrophe', 'earthquake', 'cholera epidemic', 'flood', 'suffering', 'assault', 'ordeal', 'explosion', 'famine', 'trauma', respectively.

¹⁰ 'secret', 'treasure', 'topic', 'enigma', 'question mark', 'dilemma', 'mystery', 'question', 'lesson', respectively.

the example sentences—despite the fact that such forms are commonly included in explanatory dictionaries.

6. Conclusion

Summarizing, based on our multi-phase evaluation, we believe that the proposed method—organizing adjective senses around co-occurring noun groups and generating definitions via structured prompts—is a viable and effective approach for building the microstructure of a monolingual dictionary. ChatGPT proved to be a useful tool in this process, particularly when provided with precise instructions and controlled inputs, as it was able to generate contextually appropriate, lexicographically relevant definitions and example sentences.

However, the output quality remained inconsistent across prompts and languages, and issues such as hallucinated meanings or nominal-style interpretations were not uncommon. These limitations underline that, while large language models can substantially support lexicographic work, thorough human post-editing remains indispensable to ensure linguistic accuracy, terminological precision, and structural coherence.

LLMs cannot replace human lexicographers, but our results demonstrate that, when guided by structured prompts derived from graph-based semantic models, they can serve as reliable assistants in the compilation of high-quality monolingual dictionaries.

7. References

- Ah-Pine, J., & Jacquet, G. (2009). Clique-based clustering for improving named entity recognition systems. In A. Lascarides et al. (Eds.), *Proceedings of the 12th Conference of the European Chapter of the ACL (EACL 2009)* (pp. 51–59).
- Apresjan, J. D. (1974). Regular polysemy. *Linguistics*, 12(142), pp. 5–32.
- Haber, J., & Poesio, M. (2024). Polysemy—Evidence from linguistics, behavioral science, and contextualized language models. *Computational Linguistics*, 50(1), pp. 351–417. https://doi.org/10.1162/coli_a_00500
- Comprehensive Dictionary of Hungarian = Ittész, N. (chief ed.) (2006–). *A magyar nyelv nagyszótára*, I–VIII. Budapest: MTA Nyelvtudományi Intézet.
- Concise Dictionary of Hungarian = Pusztai, F. (ed.). (2003). *Magyar értelmező kéziszótár*, Budapest: Akadémiai Kiadó.
- Héja, E., & Ligeti-Nagy, N. (2022a). A proof-of-concept meaning discrimination experiment to compile a word-in-context dataset for adjectives—A graph-based distributional approach. *Acta Linguistica Academica*, 69(4), pp. 521–548.
- Héja, E., & Ligeti-Nagy, N. (2022b). A clique-based graphical approach to detect interpretable adjectival senses in Hungarian. In D. Ustalov, Y. Gao, A. Panchenko, M. Valentino, M. Thayaparan, T. H. Nguyen, G. Penn, A. Ramesh, & A. Jana (Eds.), *Proceedings of TextGraphs-16: Graph-based methods for natural language processing*, pp. 35–43. ACL.

- Héja, E., Ligeti-Nagy, N., Simon, L., & Lipp, V. (2023). An unsupervised approach to characterize the adjectival microstructure in a Hungarian monolingual explanatory dictionary. In M. Medved, M. Měchura, C. Tiberius, I. Kosem, J. Kallas, M. Jakubíček, & S. Krek (Eds.), *Electronic lexicography in the 21st century (eLex 2023): Invisible lexicography. Proceedings of the eLex 2023 conference*, pp. 150–167. Lexical Computing.
- Héja, E., Gábor, K., Simon, L., & Lipp, V. (2024a). Graph-based detection of Hungarian adjectival meaning structures via monolingual static embeddings. In Š. Despot, K. Ostroški Anić, & I. Brač (Eds.), *Lexicography and Semantics: Proceedings of the XXI EURALEX International Congress*, pp. 235–247. Zagreb: Institute of Croatian Language and Linguistics.
- Héja, E., Gábor, K., Györffy, A., Ligeti-Nagy, N., Simon, L., & Lipp, V. (2024b). Melléknevek disztribúciós és szemantikai mintázatai. In V. Lipp, N. Ligeti-Nagy, & L. Simon (Eds.), *Prószéky Gábor 70: PG70 – Ünnepi kötet*, pp. 44–51. Budapest: HUN-REN Nyelvtudományi Kutatóközpont.
- Jakubíček, M., & Rundell, M. (2023). The end of lexicography? Can ChatGPT outperform current tools for post-editing lexicography. In M. Medved, M. Měchura, C. Tiberius, I. Kosem, J. Kallas, M. Jakubíček, & S. Krek (Eds.), *Electronic lexicography in the 21st century: Invisible lexicography (eLex 2023): Invisible lexicography. Proceedings of the eLex 2023 conference*, pp. 518–533. Lexical Computing.
- Ježek, E. (2016). *The lexicon: An introduction*. Oxford University Press.
- Kiefer, F. (2000). *Jelentélmélet*. Budapest: Corvina.
- Kiefer, F. (2008). A melléknevek szótári ábrázolásáról. In *Strukturális magyar nyelvtan 4: A szótár szerkezete*. Budapest: Akadémiai Kiadó, pp. 505–538.
- Kilgarriff, A., Husák, M., McAdam, K., Rundell, M., & Rychlý, P. (2008). GDEX: Automatically finding good dictionary examples in a corpus. *Proceedings of the XIII EURALEX international congress*, Barcelona, pp. 425–432.
- Miller, G. A. (1995). WordNet: A lexical database for English. *Communications of the ACM*, 38(11), pp. 39–41.
- Ostermann, Carolin. (2015). *Cognitive Lexicography: A New Approach to Lexicography Making Use of Cognitive Semantics*, Berlin, München, Boston: De Gruyter. <https://doi.org/10.1515/9783110424164>
- Rehurek, R., & Sojka, P. (2011). *Gensim—Python framework for vector space modeling*. NLP Centre, Faculty of Informatics, Masaryk University, 3(2).
- Štrkalj Despot, K., Ostroški, A., Brač, I. (eds.) (2024). *Lexicography and Semantics: Proceedings of the XXI EURALEX International Congress 8–12 October 2024*. Zagreb: Institute of Croatian Language and Linguistics.
- Svensén, Bo. (2009). *A Handbook of Lexicography: The Theory and Practice of Dictionary-Making*, Cambridge: Cambridge University Press.
- Zinoviev, D. (2018). *Complex network analysis in Python: Recognize - Construct - Visualize - Analyze - Interpret*. The Pragmatic Bookshelf.

This work is licensed under the Creative Commons Attribution ShareAlike 4.0 International License.

<http://creativecommons.org/licenses/by-sa/4.0/>

