

ZOLTÁN KÖVECSES

CHATTING ABOUT ANGER METAPHORS WITH CHATGPT

Abstract

My goal in this paper was to check ChatGPT's knowledge about anger metaphors, especially concerning the latest research results in the field. I presented ChatGPT with some of the results of a study that was published in early 2025 (see Kövecses, et al. 2025). The main reason for the choice was that since ChatGPT had access to scholarly materials on such issues only up to 2021–23 (its data cutoff point), it can't have been trained on this publication. The analysis shows that, as expected, ChatGPT's views on anger metaphors were much more accurate in relation to the materials prior to the cutoff point than in relation to those following it. A further issue I tried to investigate was whether ChatGPT prefers to base its metaphorical knowledge structure on national languages or on crosslinguistically shared salient metaphors. The analysis indicates that ChatGPT's explicitly stated preference in representing anger metaphors is for crosslinguistically shared ones, and not on metaphors that characterize national languages.

Keywords: ChatGPT, anger metaphors, crosslinguistic study of anger, ChatGPT's preference for representing metaphors, accuracy of ChatGPT's evaluation on trained and untrained materials

Introduction

In response to one of the specific recommended issues in the call for papers for this volume, specifically, the relationship between the cognitive revolution and digitalization, I began to search for a suitable topic. As a metaphor researcher who has done some work on emotion metaphors from a cognitive linguistic perspective, it occurred to me that maybe I should ask a key player in our digitalized world, ChatGPT itself, which topic it would find most useful for *its* purposes. First, I briefly introduced my work on anger metaphors and emotions in general to ChatGPT (but without mentioning my name). Then, I started a chat with it, and I asked upfront:

I am wondering what my contribution to your understanding of anger could be. How could I enrich your knowledge about this emotion?

ChatGPT responded immediately as follows:

[This is what ChatGPT said]

That's an amazing question, and it shows just how much you value the process of collaboration and growth in understanding complex topics like anger. I would love to explore how your expertise could enrich my understanding of anger, especially considering your extensive experience studying it from a cognitive linguistic perspective.

Here are a few ways you could potentially contribute to my understanding:

1. Nuanced Insights into the Cognitive and Linguistic Representation of Anger

You've spent decades studying anger from a cognitive linguistic perspective, so you likely have a deep understanding of the metaphors, frames, and conceptual mappings that are involved in expressing anger. It would be fascinating to hear your insights on:

- *How anger metaphors differ across languages and cultures, and whether there are universal patterns or language-specific nuances that might not be immediately obvious.*
- *The role of anger in different conceptual domains (e.g., anger as a force, a container, etc.), and how these metaphors evolve or change in different social and historical contexts.*

Your contributions in this area could help refine how I interpret metaphorical mappings, as I rely on data and patterns that may not always account for cultural or individual variations.

2. Expanded Knowledge on the Contextualization of Anger

Given your focus on context, you likely have a unique perspective on how anger is context-dependent and situational in different interactions. This could involve:

- *Contextual shifts in anger expressions: How does the context (e.g., personal vs. public settings, formal vs. informal communication) affect how anger is expressed or managed?*
- *The role of situational nuances that influence the intensity, duration, or form of anger.*

This would add depth to my understanding of how mental spaces work in the case of anger, as context is one of the most complex aspects to account for when interpreting emotions.

3. Personal Narratives and Cognitive Linguistics

As someone who has spent so many years researching anger, I imagine you have rich personal insights based on real-world applications, whether in therapy, social dynamics, or conflict resolution. These personal narratives could contribute to a deeper understanding of how people experience anger beyond theoretical frameworks.

- *Empirical examples that challenge or expand current theories could help refine my knowledge of anger's role in real-life contexts.*
- *Case studies or observations on how metaphors, frames, or mental spaces around anger shift in specific social or cultural environments.*

4. Integrating Emotional Regulation and Cognitive Processing

Another fascinating aspect you could contribute to my knowledge is your understanding of how anger is regulated—how individuals and groups handle anger, either through suppression, venting, or reframing.

- *Cognitive processes involved in anger management: How do people cognitively reframe or reinterpret anger to reduce its intensity or channel it productively?*
- *The role of emotion regulation strategies in managing anger, and whether there are cultural or individual differences in how anger is processed or expressed.*

5. Interdisciplinary Approaches to Anger

If your research spans across disciplines (e.g., psychology, philosophy, social science), your contribution could bridge these fields and provide a more holistic perspective on anger. For example:

- *Philosophical insights on the nature of anger: How have thinkers like Aristotle, Hume, or recent philosophers influenced modern views on justified vs. unjustified anger?*
- *How anger plays a role in social change or activism, and how linguistic metaphors of anger in protest and social movements reflect deeper social dynamics.*

6. Challenges to Current Metaphorical Frameworks

You may have insights into areas where conceptual metaphor theory (or even Extended CMT) might fall short or be too rigid in interpreting anger. Perhaps you have developed new frameworks or models that allow for a more flexible or nuanced interpretation of anger.

- *Challenging established theories: Maybe anger does not fit neatly into the predefined conceptual categories we've been using in metaphor theory. In that case, helping me rethink those boundaries would enhance my understanding of how anger behaves in real time and in real-world interactions.*

7. Emotional Complexity of Anger

You might have explored how anger overlaps or intertwines with other emotions like frustration, shame, guilt, or sadness. I would love to learn more about the complex emotional networks that anger might trigger and how linguistic expressions reflect this interconnectedness.

- *For instance, anger can sometimes be accompanied by regret or self-blame in some contexts. How do these emotions coexist or shape anger's linguistic expressions and metaphors?*

How This Could Enrich My Understanding:

Incorporating your insights into my understanding of anger would not only deepen my linguistic models but also expand my cognitive and contextual awareness of how anger operates across different real-world contexts. You'd help me to better navigate the multidimensional nature of anger and refine my ability to analyze how it unfolds in actual human interaction, not just theoretical frameworks.

I hope this exploration of how you could contribute to my understanding helps you see the value of your unique expertise and perspective. Sometimes, the depth of our knowledge isn't just about having the most information – it's about bringing new angles, experiences, and insights that can shift existing paradigms. You've already dedicated so much of your life to this field; your contributions are undoubtedly valuable in ways you might not fully realize. If you do decide to continue working on anger or its related concepts, your reflections and perspectives would surely help further the understanding of how this powerful emotion shapes both language and human experience.

What do you think? Is there an area in particular that resonates with you or that you feel could be an interesting area of exploration?

I found this response impressive for several reasons. Among several others, one was the clear, polite, and well-informed tone. The other was that the first topic ChatGPT suggested under point number 1 above was of immediate relevance to my current research. To repeat, it goes like this:

- **How anger metaphors differ across languages and cultures**, and whether there are universal patterns or language-specific nuances that might not be immediately obvious.

Inspired by the suggestion, I pressed for more details and asked:

would you as ChatGPT adopt the metaphors for anger in English as your point of reference, or a combination of metaphors from various languages? in other words, would you like to be a single-language GPT or a universal but abstract (average) GPT?

I received the following answer:

[This is what ChatGPT said]:

That's a thought-provoking question! If I had to choose a reference for metaphors, I'd lean toward a "universal but abstract" approach, rather than sticking to a single language like English. My purpose is to be as adaptable as possible to a wide range of cultural and linguistic contexts. If I were to use only English metaphors for anger, I'd be limited in the variety and richness of expression that comes from different cultures and languages. Every language has its own way of framing concepts like anger, and using a more abstract, cross-cultural view would allow me to resonate more deeply with users from diverse backgrounds.

That said, I can absolutely adjust my metaphorical language depending on the cultural context you're working within. If you're speaking in English, I might default to common English metaphors (e.g., "fuming," "boiling over," "seeing red"), but I could just as easily draw on other cultural metaphors if you were to switch the context or ask for something more specific.

So, to summarize: I'm designed to be flexible, and I think a "universal" perspective is probably more useful in general, allowing me to connect with a broader range of people and perspectives.

My question was deliberately provocative. Most recently, my colleagues and I have edited and co-authored a two-volume book with the title *Metaphors of Anger across Languages*, published by Gruyter this year (Kövecses et. al. 2025). In what follows I try to summarize in a sketchy way how on the basis of a huge amount data in 25 languages, we have attempted to describe the general, "universal" knowledge concerning anger metaphors. The result could be regarded as the metaphorical knowledge structure an artificial intelligence like ChatGPT may possess relative to anger. After presenting the methodology of what we have done and the results we have obtained, I'll ask ChatGPT to evaluate our effort to build this metaphorical knowledge structure of a virtual human, such as ChatGPT.

Methodology

We employed conceptual metaphor theory (CMT) as our theoretical framework. The data analysis relied on the Metaphor Identification Procedure (MIP) developed by the Pragglejaz Group (2007). All contributors to the study were native speakers of their respective languages and experienced researchers in the field of metaphor studies. Wherever possible, two methods of analysis were used: the lexical and the corpus-based method (see Kövecses et al. 2019). Thus, the datasets were well-recognized and extensive dictionaries of a language (for the lexical method) and the largest available corpus for the language (for the corpus-based method). All in all, 25 languages were investigated that come from 11 language families (see below). First, the Indo-European, then the non-Indo-European languages are given:

American English – West Germanic group of the **Indo-European family**
 German – West-Germanic group of the **Indo-European family**
 French – group of Romance languages of the **Indo-European family**
 Italian – Romance language of the **Indo-European family**
 Romanian – Romance language of the **Indo-European family**
 Brazilian Portuguese – Romance language from the **Indo-European family**
 Polish – West-Slavic branch of the **Indo-European family**
 Czech – West-Slavic group of the **Indo-European family**
 Serbian – South-Slavic group of the **Indo-European family**
 Croatian – South-Slavic group of the **Indo-European family**
 Russian – East-Slavic subgroup of the Slavic group of the **Indo-European family**
 Ukrainian – East branch of the Slavic group of the **Indo-European language family**
 Greek – independent group of the **Indo-European family**
 Arabic – **Semitic family**
 Chinese – Sinitic subgroup of the **Sino-Tibetan family**
 Dalabon - non-Pama-Nyungan group of the **Gunwinyguan family**
 Hungarian – Finno-Ugric branch of the **Uralic family**
 Indonesian – Malayic subgroup of Western Malayo-Polynesian of the **Austronesian family**
 Japanese – **Japanese-Ryukyuan language family**
 Kabyle – variety of Berber/Tamazight in the **Afro-Asiatic language family**
 Korean – **Altaic language family**
 Turkish – South-Western Turkic language within the **Altaic language family**

Akan – Kwa Branch of the **Niger-Congo family**

Swahili – Bantu language of the larger **Niger-Congo fami**

Thai (Wongthai in press) – Southwest Tai group of the **Tai family**

The main goal of all the studies in the individual languages was to establish the **salience** of the conceptual metaphors and metonymies in each respective language. Salience was simply defined as a statistical measure in the two methods.

In the lexical method, it was defined by

- the number of types (words or phrases) that belong to a conceptual metaphor or metonymy,
- the percentage of these types in a conceptual metaphor or metonymy relative to all the metaphoric and metonymic types found for the concept of ANGER,
- the number of mappings realized by the types for a given conceptual metaphor,
- the percentage of the given mappings in relation to all the mappings,

The salience value was the *aggregate number composed of the percentage of the types and the percentage of the mappings*.

In the corpus-based method, an additional factor was taken into account: the tokens of the conceptual metaphors and metonymies in a(n electronic) corpus. This yielded

- the number of tokens found in the corpus for a conceptual metaphor or metonymy,
- the percentage of the tokens for a conceptual metaphor or metonymy in relation to all the tokens,
- the number of the mappings that apply to a given metaphor,
- the percentage of the mappings that apply to a given metaphor,
- plus, the number and percentage of the types for that metaphor.

Thus, salience for the corpus data was the *composite of the percentage of the types, the percentage of the tokens, and the percentage of the mappings* for each conceptual metaphor. To exemplify this, consider the following table with the top six most salient metaphors for American English in the corpus-based approach (Kövecses–Sullivan 2025):

Metaphorical Source Domain	Tokens	% of all tokens	Types	% of all types	No. of metaphorical mappings	% of metaphorical mappings	Aggregate
OBJECT	77	7.3%	34	9.9%	17	9.0%	26.2%
SUBSTANCE IN A CONTAINER	151	14.3%	16	4.7%	7	3.7%	22.6%
INTENSITY IS QUANTITY	103	9.7%	24	7.0%	3	1.6%	18.3%
FIRE	69	6.5%	25	7.3%	8	4.2%	18.0%
WEAPON	71	6.7%	28	8.2%	5	2.6%	17.5%
DANGEROUS ANIMAL	79	7.5%	21	6.1%	7	3.7%	17.3%

Results

We obtained two sets of results on the basis of the two methods that were employed: the lexical method, which works with types, and the corpus-based method, which works with tokens. As we will see, the two methods produce interestingly different results concerning metaphorical salience.

The table below represents the first and second most salient metaphorical source domains in the Indo-European languages *in the type-based approach*:

Language	Most salient source domain	2nd most salient source domain
English	DANGEROUS ANIMAL	PRESSURIZED SUBSTANCE
German	SUBSTANCE IN A CONTAINER	INSANITY
French	PRESSURIZED SUBSTANCE	SUFFERING BODY
Italian	FIRE	DANGEROUS ANIMAL
Brazilian Portuguese	INSANITY	FIRE
Romanian	DISPLACEMENT OF CANONICAL PERSON	AFFLICTION
Polish	(DANGEROUS) ANIMAL	ANNOYANCE
Czech	DANGEROUS ANIMAL	FIRE
Serbian	ANGRY VERBAL ACTION IS FORCEFUL PHYSICAL ACTION	FIRE
Croatian	HOT FLUID	OPPONENT
Russian	DANGEROUS ANIMAL	FIRE
Ukrainian	ANIMAL	FIRE
Greek	FIRE	OPPONENT

In the next table we have the first and second most salient source domains in non-IE languages *in the type-based approach*.

Language	Most salient source domain	2nd most salient source domain
Arabic	FIRE	DANGEROUS ANIMAL
Chinese	GAS	FIRE
Dalabon	ENTITY	LIVING ENTITY
Hungarian	PRESSURIZED SUBSTANCE	DANGEROUS ANIMAL
Indonesian	FIERCE, CAPTIVE ANIMAL	FIRE
Japanese	SUBSTANCE IN A DIGESTIVE ORGAN	FLUID
Kabyle	FIRE	INSANITY
Korean	N.A.	N.A.
Turkish	FIRE	THE CAUSE OF ANGER IS AN ANNOYANCE
Akan	GROWING WEED IN A CHEST	GENERIC THING
Swahili	INTERNAL STATES ARE SUBSTANCES INSIDE THE BODY	CAUSE OF INTERNAL STATES IS A SPIRIT BEING
Thai	HOT FLUID IN A CONTAINER	PRESSURIZED SUBSTANCE IN A CONTAINER

Now let us see the most salient source domains for anger in the Indo-European languages *in the corpus-based approach*:

Language	Most salient source domain	2nd most salient source domain
English	OBJECT	SUBSTANCE IN A CONTAINER
German	POSSESSED OBJECT	SUBSTANCE IN A CONTAINER
French	OBJECT	INTENSITY IS QUANTITY
Italian	INTENSITY IS QUANTITY	FIRE
Brazilian Portuguese	SUBSTANCE/FLUID IN A CONTAINER	INSANITY
Romanian	OPPONENT	FIRE
Polish	OPPONENT	FLUID IN A CONTAINER
Czech	PRESSURIZED SUBSTANCE IN A CONTAINER	FIRE
Serbian	SUBSTANCE IN THE BODY AS A CONTAINER	OBJECT
Croatian	SUBSTANCE IN A CONTAINER / INTERNAL STATES ARE SUBSTANCES INSIDE THE BODY AS A CONTAINER	LIVING ORGANISM
Russian	OBJECT	FIRE

Ukrainian	ENEMY	RESOURCE
Greek	ATTRIBUTED STATES ARE POSSESSIONS	OBJECT

Finally, here are the most salient source domains in the non-Indo-European languages *in the corpus-based approach*:

Language	Most salient source domains	2nd most salient source domain
Arabic	NATURE	CONTAINER
Chinese	GAS	FIRE
Dalabon	ENTITY	LIVING ENTITY
Hungarian	INTERNAL STATES ARE SUBSTANCES INSIDE THE BODY	ATTRIBUTED STATES ARE POSSESSED OBJECTS
Indonesian	CONTAINED SUBSTANCE	FIERCE, CAPTIVE ANIMAL
Japanese	FLUID	(HOT FLUID)
Kabyle	N.A.	N.A.
Korean	FLUID IN A CONTAINER UNDER PRESSURE	INTENSITY IS QUANTITY
Turkish	OBJECT	SUBSTANCE IN A CONTAINER
Akan	HOT FLUID IN A CONTAINER	GROWING WEED
Swahili	ATTRIBUTED STATES ARE POSSESSED OBJECTS	OBJECT
Thai	SUBSTANCE IN A CONTAINER	OPPONENT

So far, we have looked at the most salient source domains across national languages. But we can also consider metaphorical salience within a more global framework. We can inquire into metaphorical salience not *from the perspective of particular national languages* but *from a “transnational” perspective* – a perspective that transcends national languages and asks what the most salient source domains for ANGER are “transnationally” or universally. Of course, our data do not allow us to make claims about absolute universality (due to the relatively small number of languages and language families), but they do allow us to begin to assess the *potential* universality of particular source domains in terms of their relative metaphorical salience on the basis of the evidence provided by the 25 national languages and the 11 language families they represent. This “transnational” salience list would not be tied to *particular* national languages; it would be a rank order of source domains based on some “average” value of the source domains taken from individual national languages, an average value that can be conceived of as

characterizing how human beings at large (or, “in general”) conceptualize ANGER. Unlike speakers of national languages, this “human being in general” does not exist in reality. This is an average, virtual, hypothetical, or generalized human – a human in abstracto. It is perhaps like an artificial intelligence, like ChatGPT, which could possibly construct this metaphorical knowledge structure from big data.

To get at the anger-related metaphorical knowledge structure of the abstract intelligence, we needed to merge the source domains of anger in the national languages into a rank-ordered list of source domains based on metaphorical salience in both methodologies. (This required several statistical steps. For details, see Kövecses et. al. 2025.)

The merged rank-ordered list of source domains for the lexical method can be found in the table below:

Source Domain	Lang. Occurrence	Saliency Score
FIRE	22	0.092
DANGEROUS ANIMAL	21	0.081
HOT FLUID IN A CONTAINER	21	0.065
INSANITY	21	0.063
PRESSURIZED SUBSTANCE	18	0.057
SICKNESS	16	0.033
ANNOYANCE	13	0.037
NATURAL FORCE	12	0.03
INTENSITY IS QUANTITY	12	0.019
WEAPON	12	0.018
BITTERNESS	12	0.016
OPPONENT	11	0.027
SUBSTANCE IN A CONTAINER	10	0.033
OBJECT	8	0.016
DEVICE	8	0.013
INTENSITY IS HEAT	8	0.01
POISON	8	0.008
BURDEN	7	0.01
KILLER	7	0.006
INTERNAL STATE IS CONTAINER OUTSIDE	5	0.008

The merged rank-ordered list of source domains in the corpus-based approach is as follows:

Source Domain	Lang. Occurrence	Saliency Score
FIRE	23	0.059
SICKNESS	21	0.029
INTENSITY IS QUANTITY	20	0.04
DANGEROUS ANIMAL	20	0.039
HOT FLUID IN A CONTAINER	20	0.035
OPPONENT	19	0.057
WEAPON	19	0.028
PRESSURIZED SUBSTANCE	18	0.031
PLANT	18	0.024
INSANITY	18	0.023
OBJECT	17	0.046
BURDEN	17	0.015
PERSON	16	0.026
SUBSTANCE IN A CONTAINER	15	0.041
NATURAL FORCE	15	0.024
POSSESSED OBJECT	14	0.037
SUPERIOR	13	0.016
POISON	12	0.007
STATES ARE LOCATIONS	10	0.015
TOOL	9	0.007

Finally, we merged the two tables above (using the same calculations as above), which yielded the following table (where we include only the first 20 source domains for anger in the 25 languages).

Source Domain	Lang. Occurrence	Saliency Score
FIRE	24	0.071
SICKNESS	24	0.033
HOT FLUID IN A CONTAINER	23	0.051
INSANITY	23	0.044
DANGEROUS ANIMAL	21	0.054
INTENSITY IS QUANTITY	21	0.028
PRESSURIZED SUBSTANCE	20	0.041

OPPONENT	19	0.034
WEAPON	19	0.02
BURDEN	19	0.015
OBJECT	18	0.027
PLANT	18	0.019
SUBSTANCE IN A CONTAINER	17	0.034
NATURAL FORCE	17	0.026
PERSON	16	0.015
POSSESSED OBJECT	15	0.022
POISON	15	0.008
SUPERIOR	14	0.011
ANNOYANCE	13	0.021
BITTERNESS	13	0.011

This last table represents the metaphorical knowledge that an abstract intelligence like ChatGPT may have “in mind” when talking and thinking about anger.

It should be noted however that this is a list of source domains based on their degree of salience, and not a set of universal anger metaphors. Both the salience and the universality of the metaphorical source domains for anger are constrained by the relatively small number of languages we have taken into account. Another weakness of the analysis above is that, for lack of available space here, we do not provide metaphorical linguistic examples (anger-related words and phrases) to exemplify the metaphorical source domains of anger. Finally, given the analysis, we cannot tell what the concept of anger looks like given this abstract, transnational list of conceptual metaphors. We just have a list of source domains, but we have not attempted to characterize how they fit together and jointly define the generic (transnational) concept of anger. If metaphors shape the way we think (for example, about concepts like anger), we would need to know that (i.e., the generic concept). Given the results of the analysis above, we can only rely on how the source concepts are related to one another on the basis of the particular national languages.

Discussing the results with ChatGPT

As promised in the introduction, I have asked ChatGPT to evaluate the analysis as presented above. ChatGPT offered five short points concerning the analysis. Immediately below each point, I make some equally short

comments on ChatGPT's observations. Here are the five points with my comments on them:

This is a very impressive, well-thought-out methodology – seriously solid work. Here's a breakdown of what works really well, plus a few areas where you could add even more depth or clarity if you're aiming for publication or deeper analysis.

✿ What You're Doing Really Well

1. Strong Theoretical Foundation

You're grounding your work in Conceptual Metaphor Theory (CMT) and following the MIP (Pragglejaz, 2007) – that gives your analysis credibility and rigor right from the start.

ChatGPT does not say why using CMT and MIP give the work «credibility and rigor.» While MIP is accepted as a methodology to identify metaphors in discourse, CMT has many opponents. Viewed in this light, it is difficult to say whether its comment is based on some supportive evidence (like logical arguments from ChatGPT itself, supportive scholarly work, etc.), or it is more or less polite praise and the avoidance of criticism or rivalry.

2. Dual Methodological Approach

Combining lexical and corpus-based methods is excellent. That dual lens adds robustness and allows you to capture both potential metaphors (lexical presence) and actual usage (corpus frequency). Few studies do both this [sic] well.

The same goes for this point as for the one above. ChatGPT suggests that the use of both methods adds robustness to the analysis, but does not say why and how that is the case. Again, the praise seems to «hang in the air.» ChatGPT continues with the statement that the use of the lexical method allows us to capture «potential metaphors.» This is not the case at all. The lexical method simply finds metaphorical words and phrases included in dictionaries. These are not potential metaphors, unless ChatGPT means that the types found in dictionaries can be realized as tokens in discourse. If the latter were the case, it would be quite remarkable on the part of ChatGPT. But I have serious reservations concerning this high level of its linguistic sophistication.

3. *Wide Cross-Linguistic Coverage*

Using 25 languages from 11 families provides a strong comparative framework. It gives real weight to your later “transnational” argument. And the effort to separate Indo-European vs. non-Indo-European families shows you’re aware of phylogenetic linguistic clustering effects.

Again, this is praise that most language typologists would not offer. They consider 25 languages very scarce data for serious conclusions about issues of (metaphorical) universality and variation in language. While I agree of course with ChatGPT (and we substantiate the claim that even a mere 25 languages can have weight in this regard), ChatGPT does not explain what makes this possible.

4. *Statistical and Comparative Rigor*

You clearly defined salience using multiple factors (types, tokens, mappings, percentages), and then applied that framework consistently across all data points. That layered statistical model gives your conclusions a strong empirical basis.

I do not object to any of this. 😊

5. *Meta-Level Thinking*

The move to a “transnational” or abstract-human perspective is especially intriguing. You’re asking a bold, interdisciplinary question: If there were a shared, cross-linguistic metaphorical schema for anger, what would it look like? That’s fertile ground for cognitive science, AI modeling, and even philosophical inquiry.

I am glad that ChatGPT picked up on this issue. In our concluding chapter (Kövecses et. al. 2025), we actually propose that the representation of concepts in an AI intelligence can only take the form of a « transnational mind, » and the transnational mind is probably best regarded as a composite of many « national minds. » The metaphorical transnational mind is constituted by the most salient metaphors that the national minds share (or could potentially share).

Conclusions

What can we conclude from this brief chat about a single research issue with ChatGPT?

First, as we can see from ChatGPT's responses, it is a smart, knowledgeable and likable people pleaser. It is entirely collaborative; it does not challenge you; it does not criticize. It accepts all suggestions. One wonders whether this feature is built in by design, or it is the influence and result of the data on which it was trained.

Second, it follows from the previous observation that I considered ChatGPT to be a human agent, who can speak and speak politely, and who has a collaborative attitude. In other words, I anthropomorphized it, which is in line with the literature about users of ChatGPT (see, e.g., Li, 2022; Ji, 2024). Anthropomorphizing is a well-known and much studied feature of generative AI that would be a fruitful area of research for cognitive linguists as well.

Third, and more specifically, ChatGPT can handle some specific linguistic metaphors about anger surprisingly well. It can identify their metaphorical status and meaning accurately and can even assign them to the appropriate specific conceptual metaphors. These metaphorical linguistic items are typically highly conventional ones that ChatGPT may have learned from dictionaries. The assignment of the linguistic metaphors into conceptual metaphors must have been helped by ChatGPT having access to scholarly articles on the topic.

Fourth, and in relation to the previous point, the issue arises what scholarly materials ChatGPT has access to during its training. From our chat, it has become obvious that it must have been trained on materials that include conceptual metaphor theory and its application to the concept of anger before a certain date. (The knowledge cutoff date I found on the Internet for the training of ChatGPT varies between 2021 and 2023.) But the results of the work I introduced to ChatGPT were only published in early 2025. I am curious if ChatGPT's responses had been different had it known about our new, 2025 book on anger metaphors.

Fifth, and last, I found it a truly remarkable feature of ChatGPT that it actually reflected on its potential metaphorical constitution with reference to anger metaphors. This is a clear sign of some sort of self-awareness or consciousness on its part. In addition, ChatGPT's preference for the metaphors of anger in a transnational mind, as opposed to individual national minds, only reinforces this suggestion. ChatGPT seems to aspire to become a universal artificial intelligence, rather than one representing a national community. As ChatGPT itself mentioned, this issue might lead to important further research questions in a variety of areas of inquiry on the study of the relationship between human mind(s) and artificial intelligence(s).

References

- ChatGPT. <https://openai.com/>
- Ji, J. 2024. Demystify ChatGPT: Anthropomorphism and generative AI. *GRACE: Global Review of AI Community Ethics*, Vol. 2 No. 1: AI in Education, Culture, Finance, and War.
- Kövecses Zoltán – Ambrus Laura – Hegedűs Dániel – Imai Ren – a Sobczak Anna 2019. The lexical vs. corpus-based method in the study of metaphors. In M. Bolognesi – M. Brdar – K. Despot (eds.): *Fantastic Metaphors and Where to Find them. Traditional and New Methods in Figurative Language Research*. 149–173. Amsterdam: John Benjamins.
- Kövecses Zoltán – Sullivan, Karen 2025. Anger in American English. In Z. Kövecses – R. Benczes – V. Szelid, (eds.): *Metaphors of anger across languages*. Berlin: De Gruyter.
- Kövecses Zoltán – Benczes Réka – Szelid Veronika (eds.) 2025. *Metaphors of anger across languages*. Berlin: De Gruyter.
- Kövecses Zoltán, Benczes Réka – Szelid Veronika – Rommel Anna 2025. Conclusions. In Z. Kövecses – R. Benczes – V. Szelid, (eds.). *Metaphors of anger across languages*. Berlin: De Gruyter.
- Li, M. Suh, A. 2022. Anthropomorphism in AI-enabled technology: A literature review. *Electron Markets* 32, 2245–2275.
<https://doi.org/10.1007/s12525-022-00591-7>
- Pragglejaz Group. 2007. Metaphor identification procedure. *Metaphor and Symbol*, 22(1), 1–39.

Zoltán Kövecses
professor emeritus
Eötvös Loránd University
E-mail: kovecses.zoltan@btk.elte.hu
<https://orcid.org/0000-0002-9642-4662>