

Mapping Political Metaphors in Taiwanese Social Media: A Large Language Model Approach

Chang-Yu Tsai

National Institutes of Cyber Security

entro.changyu@nics.nat.gov.tw

Metaphor is at the heart of human cognition. In Taiwan, political metaphors reflect wider ideological tensions shaped by social culture, political development, and historical narratives; accordingly, this study asks which metaphorical mappings structure contemporary political discussion in Taiwanese social media. Compared with formal political texts, social media discourse is highly heterogeneous, featuring colloquialisms, informal syntax, culture-specific references, and novel context-dependent metaphors that remain difficult for computational models to identify and explain (Huang & Liu, 2026). Yet most Taiwan-focused studies have concentrated on manual analysis of formal documents, such as presidential transcripts and the Constitution (Chiang & Chiu, 2007; Hsu et al., 2022), leaving large-scale social media corpora underexplored. Because manual annotation is highly time- and labour-intensive (Fuoli et al., 2026), this study positions Large Language Models (LLMs) as reproducible methodological infrastructure for scalable analysis, while the primary contribution remains the substantive mapping of Taiwanese political metaphors.

Methodologically, this study builds a reproducible workflow for corpus filtering, LLM-assisted metaphor extraction, and quantitative comparison. We initially planned to compare four local-election years (2014, 2018, 2022, and 2026), but restricting the corpus to event-specific time windows and the HatePolitics board left only approximately 239 usable posts for 2018 and 16 for 2022. To avoid weakening cross-period validity with very small or temporally diffuse samples, the main analysis focuses on the two comparable corpora, 2014 and 2026; the larger 2026 corpus was undersampled to match the 2014 corpus at 624 posts per period. For metaphor extraction, we use Chain-of-Thought (CoT) prompting because completed annotation remains insufficient for robust fine-tuning, and we adopt the Chinese Metaphor Dataset with Annotated Grounds (CMDAG) protocol to extract tenor–vehicle–ground triples (Huang & Liu, 2026; Shao et al., 2024). These triples capture contextual vehicles and explicit grounds before being abstracted into broader conceptual source and target domains (Liu et al., 2026). All LLM calls use fixed prompt templates and temperature 0.0; the final run uses Gemini 3.1 Pro Preview with medium thinking level. A repeated-run diagnostic using a fixed 10% sample of the balanced corpus (124 posts, 62 per period) showed that the pipeline produced comparable extraction volumes (154–175 metaphor rows), broadly stable major source-domain rankings, and reasonable domain-level pairwise Jaccard similarity (0.41–0.52). This pattern supports the use of the model for identifying broad tendencies and points to the value of further human review for individual outputs. We therefore treated LLM outputs as candidate annotations and conducted a post-hoc human review of a fixed 10% sample of LLM-extracted candidate rows. The review coded cases as accepted, revised, or rejected: revision was used when a recoverable metaphor was present but the extracted triple or domain label was wrong, whereas rejection was reserved for cases where no political metaphorical mapping could be established from the evidence text. Recurrent error patterns identified in the review were then used to guide targeted correction of similar cases in the full output.

Result-wise, our analysis identifies three stable macro-level mappings in both periods: POLITICS IS ANIMAL, POLITICS IS WAR, and POLITICS IS ENTERTAINMENT. These macro-level mappings are aggregated from finer tenor–vehicle–ground triples and target–source domain pairs.

They therefore summarise the overall distribution of source domains, while still allowing the analysis to trace each pattern back to more specific political targets, such as politicians, parties, supporters, or crowds. In this workflow, the extracted *grounds* make these fine-grained mappings more inspectable and help guide human revision of extracted triples and domain labels. Further analysis of target-level differences within each source domain remains an important direction for the next stage of analysis. In this sense, LLMs substantially accelerate the work by producing structured mappings that can be inspected and revised, while expert linguistic insight is needed to design prompts that effectively guide the model and to judge the resulting mappings in the final review.

References

- Chiang, W.-y., & Chiu, S.-h. (2007). The conceptualization of state: A comparative study of metaphors in the r.o.c. (taiwan) and u.s. constitutions. *Concentric: Studies in Linguistics*, 33(1), 19–46.
- Fuoli, M., Huang, W., Littlemore, J., Turner, S., & Wilding, E. (2026). Metaphor identification using large language models: A comparison of rag, prompt engineering, and fine-tuning. *Applied Corpus Linguistics*, 6, 100204. <https://doi.org/10.1016/j.acorp.2026.100204>
- Hsu, H.-L., Lai, H.-l., & Liu, J.-S. (2022). Democracy in taiwanese presidential inaugural addresses: Metaphors, source domains, scenarios, and ideologies. *Concentric*, 48(2), 212–248. <https://doi.org/10.1075/consl.22006.hsu>
- Huang, W., & Liu, M. (2026). Interpretable chinese metaphor identification via llm-assisted mipvu rule script generation: A comparative protocol study. *arXiv preprint arXiv:2603.10784*.
- Liu, H., He, C., Zhu, S., Niu, C., & Jia, Y. (2026). Metaphor components identification with feedback-enhanced feature-driven in-context learning. *ACM Transactions on Asian and Low-Resource Language Information Processing*, 25(4), 28:1–28:30. <https://doi.org/10.1145/3801739>
- Shao, Y., Yao, X., Qu, X., Lin, C., Wang, S., Huang, W., Zhang, G., & Fu, J. (2024). CMDAG: A Chinese metaphor dataset with annotated grounds as CoT for boosting metaphor generation. In N. Calzolari, M.-Y. Kan, V. Hoste, A. Lenci, S. Sakti, & N. Xue (Eds.), *Proceedings of the 2024 joint international conference on computational linguistics, language resources and evaluation (Irec-coling 2024)* (pp. 3357–3366). ELRA; ICCL. <https://aclanthology.org/2024.Irec-main.298/>