

**From Protective Formula to Pragmatic Marker:
The Semantic Bleaching and Pragmatic Reanalysis of Vietnamese *trộm vía***

Cao Bang Trinh Tran - National Taiwan University - tran.caobangtrinh@gmail.com

Thanh Hoa Nguyen - National Taiwan University - caroljuan27@gmail.com

Thi Thach Thao Le - National Taiwan University - jessie.thachthao@gmail.com

Keywords: *trộm vía*; Vietnamese pragmatics; grammaticalization; politeness; digital discourse.

Although *trộm vía* originated as a Vietnamese apotropaic formula attached to compliments in order to protect the praised person from supernatural harm, contemporary written Vietnamese uses it far beyond this ritual context. It now functions as an anti-jinx hedge, a marker of modesty, an ironic flag, and a highly bleached discourse-pragmatic marker, making it a valuable case for examining semantic bleaching and pragmatic reanalysis in a culturally embedded expression. This study addresses two questions: (i) how *trộm vía* distributes across pragmatic function types in contemporary Vietnamese; and (ii) what this distribution reveals about its semantic bleaching and pragmatic reanalysis.

This study proposes a two-stage model that integrates three complementary perspectives on pragmatic change. Drawing on Carston's lexical pragmatics, the meaning of *trộm vía* is dynamically adjusted in context, producing ad hoc concepts beyond its encoded source meaning (Carston, 2002, 2019). Brown and Levinson's politeness theory accounts for the social motivation behind this expansion, especially how praise and self-praise can create interactional risk by sounding excessive, boastful, or inappropriate (Brown & Levinson, 1987). Hopper and Traugott's grammaticalization theory explains how these recurrent contextual meanings become conventionalized through semantic bleaching, subjectification, layering, and pragmatic reanalysis (Hopper & Traugott, 2003). Within this model, Stage 1 (Types 1–2) captures the phase in which *trộm vía* still preserves its apotropaic source meaning of “compliment-without-attracting-harm.” In Type 1, the expression reduces the perceived risk of spiritual harm in child-praise contexts; in Type 2, this protective meaning extends to positive outcomes more generally, where *trộm vía* functions as an anti-jinx hedge marking success as fortunate but fragile. Stage 2 (Types 3–5) begins when the expression is reanalyzed from ritual protection into interactional face management, and ultimately into a procedural discourse marker. In Type 3, it marks modest self-positioning; in Type 4, it extends to ironic or negative contexts; and in Type 5, it becomes almost fully bleached and decategorized, functioning like an evaluative modifier as in expressions such as *Các bà có tips nào để da dẻ luôn luôn trộm vía không ạ* (TV0158) “Ladies, any tips to keep [our] skin always *trộm vía*?”, where the form evaluates skin as healthy and desirable rather than warding against harm.

Empirically, the study is based on a 300-token corpus drawn from Vietnamese news media from 2006 to 2026 and social-media writing. The corpus includes three orthographic variants: canonical *vía* (n = 224), colloquial *día* (n = 69), and further-attenuated *zía* (n = 7). Each token was independently coded by three annotators using a five-type pragmatic taxonomy: Type 1 genuine apotropaic use, Type 2 positive anti-jinx hedge, Type 3 mixed-valence or modesty use, Type 4 ironic or negative use, and Type 5 fully bleached discourse-marker use. Non-analysable items (titles, duplicates, off-topic uses) were assigned to an exclude category. Inter-rater reliability is almost perfect, with Krippendorff's $\alpha = 0.985$ and pairwise Cohen's $\kappa = 0.982\text{--}0.986$ (Krippendorff, 2004; Hayes & Krippendorff, 2007).

Findings reveal gradient bleaching coupled with pragmatic reanalysis. (i) Distributional and orthographic evidence: Type 1 (genuine apotropaic) accounts for only 21.0% of tokens, while Type 2 alone reaches 44.7%; Stage 1 (Types 1–2) covers 65.7% of the corpus, and Stage 2 (Types 3–5) covers 26.6%. Type 1 is further graded across orthographic variants (27.7% in *vía*, 10.1% in *día*, 0% in *zía*; $\chi^2(1) = 8.91, p = 0.0028$), while Type 5 shows the inverse pattern, rising from 9.9% in *vía* to 28.6% in *zía*. (ii) Qualitative evidence: Four tokens display readings in which a Stage 2 function is dominant while a residual Stage 1 anti-jinx interpretation remains accessible, for example, *trộm día. mình lên làm được liền...* “*Trộm día. I leveled up right away...*” (TV0297), where standalone use indicates Type 5 detachment yet retains a Type 2 self-praise hedge reading. Such tokens anchor the two-stage model, documenting the transition from apotropaic protection toward face management and discourse marking. Together, these findings address the distributional question (RQ1) and reveal a directional bleaching-reanalysis pathway (RQ2).

By combining lexical pragmatics, politeness theory, grammaticalization theory, and corpus-based pragmatic annotation, this study shows that *trộm vía* follows a cline from ritual protection to conventionalized polysemy. Its older apotropaic function has not disappeared, but coexists with newer pragmatic functions involving luck, evaluation, politeness, irony, and social alignment. This makes *trộm vía* a valuable case for understanding how Vietnamese cultural concepts of praise, vulnerability, and anti-jinx protection enter broader processes of pragmatic change. More broadly, this case shows that theoretical linguistics remains essential in the AI era. As culturally embedded expressions circulate quickly through digital discourse, theoretical pragmatics remains crucial for recovering their layered profiles.

References

- Brown, P., & Levinson, S. C. (1987). *Politeness: Some universals in language usage*. Cambridge University Press.
- Carston, R. (2002). *Thoughts and utterances: The pragmatics of explicit communication*. Blackwell.
- Carston, R. (2019). Ad hoc concepts, polysemy and the lexicon. In K. Scott, B. Clark, & R. Carston (Eds.), *Relevance, pragmatics and interpretation* (pp. 150–162). Cambridge University Press.
- Hayes, A. F., & Krippendorff, K. (2007). Answering the call for a standard reliability measure for coding data. *Communication Methods and Measures*, 1(1), 77–89.
- Hopper, P. J., & Traugott, E. C. (2003). *Grammaticalization* (2nd ed.). Cambridge University Press.
- Krippendorff, K. (2004). Reliability in content analysis: Some common misconceptions and recommendations. *Human Communication Research*, 30(3), 411–433.