

A Rule-Based NLP Model for Hindi Noun–Adjective Grouping

Alok Kumar Mishra¹, Dr. Indra Kumar Pandey², Dr. Ambreesh Tripathi³, Sunil Shrama⁴

¹Department of Linguistics, Researcher, M. G. A. H. V, Wardha)

²Faculty of Computer Applications, Assistant Professor, Invertis University, Bareilly

³Department of Hindi/Linguistics, Assistant Professor, Enternal University, Himanchal

⁴Faculty of Computer Application, Assistant Professor, S.R.M.S., Bareilly

doi.org/10.64643/IJIRTV12I11-207095-459

Abstract—Analyzing Hindi, a morphologically rich Indo-Aryan language, presents significant challenges in Natural Language Processing (NLP), particularly in syntactic agreement, morphological variation, and semantic interpretation. For instance, adjective–noun agreement in phrases such as “अच्छा लड़का” and “अच्छी लड़की” demonstrates gender-based variation, while flexible structures like “सुंदर लड़की” and “लड़की सुंदर है” highlight syntactic diversity. This paper proposes a research-oriented rule-based model for grouping Hindi nouns and adjectives based on semantic and grammatical classification. The framework integrates linguistic principles and systematically categorizes nouns into tangible (“लाल फूल”), intangible (“सच्चा प्रेम”), animate (“खुश बच्चा”), and inanimate (“भारी पत्थर”) classes, aligning adjectives accordingly. The model processes input sentences through structured steps including tokenization, part-of-speech identification, semantic classification, and agreement matching. For example, in the sentence “सुंदर छोटी लड़की गाना गा रही है”, the system identifies and groups “सुंदर छोटी लड़की” as a coherent adjective–noun unit. An extensive example-driven evaluation demonstrates that the proposed system performs effectively in low-resource environments, where large annotated datasets are not readily available. Overall, the study contributes toward the development of explainable NLP models for Hindi and provides a foundation for future hybrid systems that combine rule-based and data-driven approaches.

Index Terms—Natural Hindi NLP, Semantic Classification, Rule-Based NLP, Adjective Agreement, Low-Resource Languages, Paninian Grammar

I. INTRODUCTION

Natural Language Processing (NLP) has significantly advanced with the emergence of deep learning and

transformer-based models; however, low-resource languages such as Hindi still require interpretable and rule-based approaches for effective linguistic analysis. Hindi, being a morphologically rich Indo-Aryan language, exhibits complex features such as gender agreement, inflectional variation, and flexible word order, which make computational processing and syntactic analysis challenging.

One of the key difficulties lies in handling agreement and structural variation. For instance:

“सुंदर लड़की” /sundər ləɽki:/ “लड़की सुंदर है” /ləɽki: suŋdər hɛ:/

Both sentences are grammatically correct and convey similar semantic meaning, yet they differ in syntactic structure. Similarly, gender-based agreement can be observed in examples such as:

“अच्छा लड़का” /əɽʃːhɑ: ləɽka:/ “अच्छी लड़की” /əɽʃːhi: ləɽki:/

Here, the adjective changes its morphological form according to the gender of the noun, highlighting the importance of agreement modeling in NLP systems.

In addition, Hindi allows multiple adjectives to modify a single noun, as seen in:

“बहुत सुंदर छोटा घर” /bəhʊt suŋdər tʃʰoːtɑ: gʱər/

Such constructions require hierarchical grouping and pose challenges for computational parsing.

Previous research in Hindi NLP has primarily focused on part-of-speech tagging, chunking, and dependency parsing. However, limited work has explored systematic noun–adjective grouping with semantic classification. For example:

“लाल फूल” /la:l pʰu:l/ (tangible)

“सच्चा प्रेम” /səɽʃːɑ: pre:m/ (intangible)

“खुश बच्चा” /kʰʊʃ bəɽʃːɑ:/ (animate)

“भारी पत्थर” /b^ha:ri: pəṭṭ^hər/ (inanimate)

These examples illustrate the necessity of integrating semantic classification with grammatical analysis. To address this gap, this paper proposes a structured rule-based model that combines linguistic rules with semantic categorization for accurate noun–adjective grouping. The approach emphasizes interpretability and linguistic transparency, making it particularly suitable for low-resource scenarios where annotated datasets are limited.

II. RELATED WORK

Grammar checking systems first emerged around 1970 with the development of Writer’s Workbench, which focused on punctuation, stylistic issues, and grammatical errors [12]. In 1981, Aspen Software Company introduced Grammatik, a pronunciation and style checker that later evolved into a full-fledged grammar checking system [14]. In 1992, grammar checking tools were commercially integrated into word processing software by companies such as WordPerfect and Microsoft, where they continue to be used in modern versions of MS Word [11].

Early grammar-checking systems were predominantly developed for the English language, with tools such as Writer’s Workbench addressing basic grammatical and stylistic errors. With time the quality of Grammatik and paid word processor checkers improved but they still depend on particular language predefined rules. Whereas tools to check grammar still is in developing stages due to the multiple word order and rich morphology. Certainly, we also got remarkable work in this field by IITs and CDAC, they have done exceptionally great work in the field of complex verb phrases construction, especially for auxiliary and conjunct verb. Out of multiple grammar checker tools in markets most of them is for English language. The research area of Hindi grammar checker is still an empty space, however IITs and CDAC is doing great job, but still there is no universally accepted procedure produced.

III. LINGUISTIC BACKGROUND

3.1 Noun Categories in Hindi

Hindi nouns can be classified semantically:
Tangible (सूत) – Physical entities

Intangible (अमूर्त) – Abstract concepts

Animate (सजीव) – Living beings

Inanimate (निर्जीव) – Non-living objects

This classification aligns with semantic ontology used in NLP systems [3].

3.2 Adjective Agreement

Adjectives in Hindi must agree with nouns in:

- Gender
- Number
- Case

Example:

अच्छा लड़का

अच्छी लड़की

Such agreement rules are critical in computational parsing [1].

IV. METHODOLOGY

4.1 System Architecture

The proposed model follows a structured rule-based pipeline designed to capture both grammatical agreement and semantic relationships between nouns and adjectives in Hindi.

1. Processing Pipeline
2. Input Sentence
3. Tokenization
4. Part-of-Speech (POS) Identification
5. Semantic Classification
6. Agreement Matching
7. Noun–Adjective Grouping

4.2 Rule-Based Processing Framework

The system is governed by linguistically motivated rules derived from Hindi grammar and semantic classification principles.

Rule 1: Adjective Identification Rule

If a word expresses a property, quality, or attribute of an entity, it is classified as an adjective.

Examples:

“सुंदर लड़की” /sundəɾ ləṛki:/

“ठंडा पानी” /ṭ^həṇḍa: pa:ni:/

“सुंदर”, “ठंडा” → Adjectives

Rule 2: Noun Identification Rule

If a word denotes an entity (person, object, concept), it is classified as a noun.

Examples:

Output:

[बहुत सुंदर छोटा] + घर

Rule 3: Coordination Rule (Compound Adjectives)

If adjectives are joined using conjunctions, all modify the same noun.

Example:

“मेहनती और बुद्धिमान छात्र”

/me:hənəti: ɔ:r bʊdd̪hima:n t̪ʰa:tr/

Output:

[मेहनती + बुद्धिमान] + छात्र

4.3 Algorithm (Formal Representation)

Input: Sentence S

Step 1: Tokenize S → W1, W2, W3...

Step 2: Identify POS for each Wi

Step 3: If Wi is noun → classify (Tangible/Intangible/Animate/Inanimate)

Step 4: If Wi is adjective → check agreement with nearest noun

Step 5: Apply grouping rules:

- Proximity rule
- Multiple adjective rule
- Coordination rule

Step 6: Output grouped structures

4.4 Example Walkthrough (End-to-End)

Input Sentence:

“एक बहुत सुंदर छोटी लड़की खेल रही है”

/Ek bəhʊt̪ sʊnd̪ər t̪ʰo:ti: ləʈki: kʰe:l rəhi: hɛ:/

Processing:

एक → Determiner

बहुत, सुंदर, छोटी → Adjectives

लड़की → Noun (Animate + Tangible)

Output:

[बहुत सुंदर छोटी] + लड़की

4.5 Algorithm

Input: Sentence S

Output: Grouped Adjective–Noun pairs

Step 1: Tokenize S into words

Step 2: Identify POS (Noun/Adjective)

Step 3: Classify noun semantically

Step 4: Check agreement rules

Step 5: Group adjective with nearest noun

V. RULE BASED FRAMEWORK

Despite its effectiveness and interpretability, the proposed rule-based model has several limitations. First, the system relies heavily on predefined linguistic rules, which makes it less adaptable to unseen or ambiguous language patterns. Hindi, being a morphologically rich and context-sensitive language, often exhibits lexical ambiguity where a word may function as both a noun and an adjective depending on context. For example, in “अच्छा”, the word may act as an adjective or a standalone expression, which cannot always be resolved using fixed rules.

Second, the model performs well on simple and moderately complex sentence structures but struggles with highly complex or compound sentences involving nested dependencies. Sentences such as “जो लड़का बहुत मेहनती और बुद्धिमान है, वह सफल होगा” require deeper syntactic parsing beyond rule-based proximity and agreement matching.

Third, the absence of statistical learning limits the system’s ability to generalize across large and diverse datasets. Unlike machine learning or transformer-based approaches, the model does not learn from data and therefore cannot automatically improve its performance over time [8], [10].

Finally, semantic classification into categories such as tangible and intangible is rule-driven and may not always capture nuanced meanings, especially in metaphorical or figurative expressions. These limitations highlight the need for integrating data-driven techniques with rule-based systems [12], [18].

VI. CONCLUSION

This paper presents a research-oriented rule-based NLP model for grouping Hindi nouns and adjectives using semantic and grammatical classification. The model systematically identifies nouns and adjectives, classifies nouns into tangible, intangible, animate, and inanimate categories, and applies agreement rules to form meaningful linguistic units.

Through extensive example-based analysis, the study demonstrates that rule-based approaches remain highly effective in low-resource language settings, where large annotated datasets are not readily available. The inclusion of linguistic principles and semantic categorization enhances the interpretability

and transparency of the system, making it suitable for educational tools, grammar checking systems, and basic machine translation applications.

Although modern NLP research is dominated by deep learning techniques, this work highlights the continued relevance of rule-based models as foundational and explainable systems, particularly for Indian languages such as Hindi [1], [4]. The proposed framework establishes a strong baseline for further research in semantic grouping and structured language processing.

VII. FUTURE WORKS

Future research can extend the proposed model in several directions. One important enhancement is the integration of machine learning and transformer-based approaches, such as BERT and multilingual language models, to improve contextual understanding and ambiguity resolution [8], [11]. A hybrid system combining rule-based and data-driven methods can leverage the strengths of both interpretability and adaptability.

Another direction involves the development of large-scale annotated corpora for Hindi noun–adjective structures, which can be used for training and evaluation of advanced models. Incorporating dependency parsing and semantic role labeling can further improve the accuracy of grouping in complex sentence structures.

Additionally, the model can be extended to handle verb phrase analysis and full sentence-level parsing, enabling broader NLP applications. Expanding the framework to other low-resource Indo-Aryan languages such as Kumaoni and Garhwali would also be a valuable contribution, supporting multilingual NLP research in the Indian context [12].

Finally, incorporating contextual embeddings and ontology-based semantic networks can enhance classification accuracy, especially in handling abstract and metaphorical expressions, thereby making the system more robust and scalable [18], [20].

REFERENCES

- [1] D. Jurafsky and J. H. Martin, *Speech and Language Processing*, 3rd ed. Hoboken, NJ, USA: Pearson, 2023.
- [2] B. Comrie, *Language Universals and Linguistic Typology: Syntax and Morphology*. Chicago, IL, USA: University of Chicago Press, 1989.
- [3] C. D. Manning and H. Schütze, *Foundations of Statistical Natural Language Processing*. Cambridge, MA, USA: MIT Press, 1999.
- [4] A. Bharati, V. Chaitanya, and R. Sangal, *Natural Language Processing: A Paninian Perspective*. New Delhi, India: PHI Learning, 1995.
- [5] K. Church and R. Mercer, “Computational linguistics using large corpora,” *Computational Linguistics*, vol. 19, no. 1, pp. 1–24, 1993.
- [6] D. Klein and C. D. Manning, “Accurate unlexicalized parsing,” in *Proc. 41st Annu. Meeting Association for Computational Linguistics (ACL)*, Sapporo, Japan, 2003, pp. 423–430.
- [7] R. Sinha, “Machine translation: An Indian perspective,” *Journal of Quantitative Linguistics*, vol. 15, no. 1, pp. 1–18, 2008.
- [8] T. Bhattacharyya, “Free word order and scrambling in Hindi,” *Linguistic Analysis*, vol. 32, nos. 1–2, pp. 33–58, 2002.
- [9] G. Leech, *Grammar for Meaning*. London, U.K.: Routledge, 2014.
- [10] B. Comrie, *Aspect*. Cambridge, U.K.: Cambridge University Press, 1976.
- [11] R. Mitkov, *The Oxford Handbook of Computational Linguistics*. Oxford, U.K.: Oxford University Press, 2005.
- [12] L. Cherry and N. MacDonald, “Writer’s Workbench,” *Bell System Technical Journal*, vol. 62, no. 6, pp. 2131–2151, 1983.
- [13] E. Brill, “A simple rule-based part-of-speech tagger,” in *Proc. Third Conf. Applied Natural Language Processing (ANLP)*, Trento, Italy, 1992, pp. 152–155.
- [14] D. Naber, “A rule-based style and grammar checker,” *Literary and Linguistic Computing*, vol. 18, no. 3, pp. 321–334, 2003.
- [15] R. Sangal, A. Sharma, and A. Bharati, “Computational grammar of Hindi,” *Language Resources and Evaluation*, vol. 41, no. 1, pp. 63–89, 2007.
- [16] Centre for Development of Advanced Computing (CDAC), *Natural Language Processing Initiatives in Indian Languages*, Technical Rep., Pune, India, 2019.