

A Study on Language Cognitive Mechanisms from the Perspective of the Integration of Psycholinguistics and AI

Lin Xin

Fuzhou Technology and Business University, Fuzhou, Fujian Province, 350715, China

Abstract: This paper explores the cognitive mechanisms of human language understanding and production from the perspective of psycholinguistics-AI integration. It reviews traditional psycholinguistic models and real-time processing mechanisms of language processing, and combines them with the technical evolution of AI language modeling, especially the development of large-scale pretrained language models driven by deep learning. Then, this paper proposes a Distributed-Symbolic Hybrid Model, which is a theoretical framework of language cognition that integrates distributed semantic representation and symbolic structure reasoning. This model is theoretically compatible with psycholinguistic evidence and AI language modeling results. By simulating the functional characteristics of semantic cortex areas (such as the TPJ) and syntactic processing related brain areas (such as Broca's area) in the human brain, the model enhances biological plausibility. The paper aims to build a language cognitive model with dynamic causal mechanisms, provide AI with interpretable cognitive constraints, and offer computational implementation tools for psycholinguistic theories.

Keywords: Language cognitive mechanism; psycholinguistics; artificial intelligence; Distributed symbolic hybrid model

1. Introduction

As artificial intelligence technologies advance rapidly, especially with breakthroughs of large-scale language models in natural language processing, the cognitive basis of language understanding has become a key interdisciplinary research topic again. Traditional psycholinguistics has accumulated a wealth of experimental data and theoretical models on language processing mechanisms, such as the modular processing hypothesis and predictive coding theory. But these theories often lack computable implementation pathways. AI language models, despite their near-human or even superhuman task performance, have a “black-box” nature, making it hard for them to reflect real human language understanding mechanisms. This has sparked debates on the nature of language understanding: is it statistical pattern matching or meaning generation? And can AI language models simulate the causal mechanisms of human language cognition? Against this backdrop, this study aims to offer a new theoretical approach. It integrates psycholinguistics, neural evidence and AI modeling methods to explore how to combine psycholinguistic theories with AI technologies and build dynamic causal models of human language understanding and production. Therefore, this paper not only focuses on theoretical construction of language cognition, but also aims to propose model frameworks with biological plausibility and computable verification potential.

2. Review of Language Cognitive Mechanisms from the Perspective of Psycholinguistics

As an important function of the cognitive system, language understanding and production have long been central to psycholinguistic research. Early studies were mostly based on the modular processing paradigm. For example, Fodor's (1983) modular view of mind suggests that language comprehension involves several relatively independent information-processing modules, such as syntactic, semantic, and phonological modules.^[1] This perspective supports the syntax-first hypothesis and emphasizes informational encapsulation and rapid automatic processing. In terms of language production, Levelt's (1989) speech production model, which divides language generation into three stages—conceptualization, formulation and articulation—and highlights the role of self-monitoring in error detection and correction, remains a classic framework.^[2] However, these models have limitations in explaining the nonlinearity and interactivity of language processing, especially in dealing with ambiguous-resolution and context-

dependent language phenomena. In recent years, the Predictive Coding Theory has become an important theory for explaining real-time language processing. Research shows that the brain continuously generates predictions during language comprehension and adjusts prediction errors based on input signals (Kuperberg, 2016).^[3] This process can be measured by event-related potential (ERP) technology; for example, the N400 component reflects semantic expectancy violation, and the P600 component corresponds to syntactic violations. These neuro-evidence with high-temporal resolution indicate that language comprehension is a highly dynamic and context-sensitive process.

Despite the continuous evolution of psycholinguistic theories, there are still two major controversies. On the one hand, the modular models emphasize informational encapsulation and bottom-up automatic processing but cannot explain phenomena like context-dependent and prediction-based language comprehension. On the other hand, the Predictive Coding Theory can describe the dynamic nature of language processing and cognitive control but does not clarify whether the brain has symbolic language-manipulation capabilities. Based on these theoretical developments, this study focuses on two core questions: whether language comprehension is driven by bottom-up stimuli or modulated by top-down cognitive control, and whether language cognition involves symbolic-manipulation capabilities or just distributed-activation states.

3. Progress and Cognitive Inspiration of AI Language Models

3.1 Traditional Symbolism-based AI and Linguistic Modeling

Early AI language modeling, influenced by rationalism and formal logic, emphasized symbol manipulation and rule-based systems. Chomsky's formal grammar systems, such as context free grammar (CFG) and transformational generative grammar (TSG), were widely used in syntactic analysis and semantic expression. Rule-based systems dominated early machine translation and question-answering systems. But it relied heavily on manually crafted rules and lacked generalization ability and semantic understanding depth. Symbolism-based AI assumed that language understanding equals symbol manipulation, which maps language to logical propositions and achieve "understanding" through formal reasoning. However, this approach struggled with language ambiguity and context dependence, showing clear limitations in natural dialogues or language variation.

3.2 The Rise of Statistical Language Models

At the end of the 20th century and the beginning of the 21st century, Statistical language models (SLMs) gradually replaced rule-based systems. N-gram probabilistic models can capture cooccurrence relationships between words, thus predicting the probability of the next word. Although SLMs enhanced the flexibility and adaptability of language models, they were limited by the local context window and couldn't effectively handle long-distance dependencies and structural language features. Probabilistic context-free grammar (PCFG), which tried to introduce syntactic structure into statistical models, still failed to solve the problem of semantic-level understanding. During this phase, AI language modeling didn't touch the cognitive essence of language understanding and was more of a probability-driven language expression optimization tool.

3.3 Breakthroughs in Deep Learning and Neural Language Models

The rise of deep learning has revolutionized language modeling. The emergence of recurrent neural networks (RNNs), long short-term memory networks (LSTMs), and transformer architectures has enabled models to handle long-distance contextual dependencies and achieve end-to-end learning of language representation. Large-scale pre-trained language models (PLMs) such as BERT and GPT have demonstrated remarkable capabilities in language understanding and generation, even rivaling human performance in many NLP tasks. These models, based on the theory of distributed semantics, embed language units into continuous vector spaces, endowing them with compositionality and generalization ability. The self-attention mechanism has become mainstream in language modeling by providing models with global context awareness.^[4]

3.4 Cognitive Analogy and Philosophical Reflections on AI Language Models

Despite the engineering success of AI language models, there is significant debate over whether they truly "understand" language. Searle's Chinese room argument suggests that a system's ability to produce

appropriate responses does not imply understanding. Similarly, while AI language models can generate coherent text, they lack genuine semantic awareness and intentionality.^[5] Furthermore, it is worth considering whether distributed semantics truly reflects human concept organization. Psycholinguistic research indicates that human language understanding is heavily dependent on background knowledge, experience, and emotional states, whereas current AI models mainly rely on statistical regularities and representational similarities. Thus, although AI language models can simulate language behavior, they fall short of the causality and contextual adaptability of human language understanding.

4. Integrated Theoretical Conception of Language Cognitive Mechanisms

4.1 Theoretical Framework: Distributed-Symbolic Hybrid Model (DSHM)

To bridge the theoretical gap between psycholinguistics and AI language models, this paper proposes a language cognitive mechanism model, integrating distributed semantics and symbolic structure reasoning, namely the Distributed-Symbolic Hybrid Model (DSHM). Based on Evans' (2003) dual-processing theory, it distinguishes two language-processing pathways. One is a fast, unconscious, distributed-activation path based on corpus statistics.^[6] The other is a slow, conscious, symbol-manipulation path relying on syntax and logic. DSHM views language comprehension as a Bayesian inference process driven by prediction-error feedback. Distributed semantics initiates activation, while syntactic-logical structures verify and redress predictions to ensure semantic coherence and task relevance. In the DSHM model, the two language-processing pathways are not independent. Instead, a dynamic regulation mechanism coordinates their synergistic effect. The distributed path quickly extracts semantic cues and forms initial understanding, suitable for high-frequency word recognition and common pattern matching. The symbolic path handles complex reasoning, ambiguity resolution, and semantic verification, crucial in tasks with deep logic or cross-modal integration.

A cognitive-control module modulates this dual-path interaction, allocating resources based on task demands, context, and experience. For simple narrative texts, the system favors the distributed path for efficiency. For abstract or logical tasks, it enhances the symbolic path to ensure coherence and logical rigor.

Computationally, DSHM combines two key AI language modeling paradigms: the global-context awareness from Transformer based self-attention and the symbolic reasoning from knowledge graphs and logic systems. It models language understanding as a prediction-error driven Bayesian process. Here, high level structures (like world knowledge, grammar rules, task objectives) generate input expectations, while low level perceptions (like word forms and sounds) provide error signals for prediction refinement. This aligns with psycholinguistic models of real time language processing and neuroscience findings on forward prediction and backward error signaling.

DSHM aims to unify different theoretical paradigms rather than replace existing theories. It accommodates psycholinguistic debates on language modularity and interactivity and offers biologically plausible cognitive constraints for AI language models.

4.2 Model Components

4.2.1 Semantic Embedding Space

Linguistic units (words, phrases, sentences) exist as high dimensional vectors in a continuous semantic space. Vector relationships reflect semantic similarity and compositionality. This space includes lexical co-occurrence information and integrates cross-modal features (e.g., visual, auditory cues), echoing the function of the brain's semantic cortex, such as the temporal parietal junction.

4.2.2 Syntax-Logic Mapper

This module transforms distributed semantics into interpretable syntactic-logical structures, enabling semantic reasoning and logical consistency. Inspired by dependency and constructional grammar in psycholinguistics, it supports recursion, categorization, and semantic role labeling, aligning with Broca's area role in syntactic processing.

4.2.3 Cognitive Control Module

This module handles task-oriented meaning selection, attention regulation, and working memory management, corresponding to the prefrontal cortex's function. It adjusts the weight of different language

processing pathways, deciding when to activate distributed intuitive processing or symbolic reasoning. This dynamic resource allocation explains individual and task-difficulty differences in language comprehension.

4.3 Dynamic Interaction Mechanism

In the DSHM model, language processing is regarded as a prediction-error feedback loop, in line with the predictive coding theory (Friston, 2005). In this mechanism, high level cognitive structures (e.g., syntax or world knowledge) generate predictions, and low level sensory inputs (e.g., word forms or sounds) provide feedback errors, thereby driving the model to update its internal representations.^[7] Additionally, the efficiency of language processing is influenced by cognitive resource allocation. Especially in complex tasks, the cognitive control module must dynamically adjust attention weights and reasoning costs to ensure language comprehension is both efficient and accurate.

4.4 Theoretical Advantages and Empirical Validation of the Model

The DSHM model shows many advantages theoretically and practically. Firstly, the DSHM model is compatible with experimental findings in psycholinguistics and achievements in AI language modeling. Its structure matches the syntactic violation (P600) and semantic anomaly (N400) responses observed in ERP studies, reflecting the immediacy and hierarchy of language processing. Secondly, it offers a unified cognitive-computational interface. Model parameters can be shared between psycholinguistic data and AI training, forming a dual way verification mechanism. Thirdly, it allows cross-modal transfer and situational adaptation in language understanding, with its built-in multimodal semantic space integrating emotions, contexts, and cultures. To enhance its practicality, the model can be trained and verified with brain computer interface (BCI) data. By capturing EEG or fMRI data during language comprehension, we can create a neural response map for language processing and use it to guide model learning. This method promises closed-loop validation and real-time intervention in language cognition.

5. Conclusion

In the perspective of psycholinguistics and AI integration, this paper proposes the Distributed-Symbolic Hybrid Model (DSHM) to unify distributed semantics and symbolic-structure reasoning. The model not only incorporates psycholinguistic insights on the real-time and predictive nature of language comprehension but also integrates cutting edge developments in AI language modeling regarding semantic representation and computational implementation. This has led to an initial realization of dynamic causal modeling in language cognition. This study, on one hand, imposes biological plausibility constraints on AI language models, propelling the development of explainable AI. On the other hand, it supplies psycholinguistic theories with computational validation tools, fostering cross-modal research into language understanding mechanisms. In the future, integrating brain computer interface (BCI) data into the model framework could achieve closed-loop validation of language-cognition mechanisms. This would further advance language science research theoretically and practically.

References

- [1] Fodor, J. A. (1983). *The Modularity of Mind: An Essay on Faculty Psychology*. MIT Press.
- [2] Levelt, W. J. M. (1989). *Speaking: From Intention to Articulation*. MIT Press.
- [3] Kuperberg, G. R. (2016). "An integrative theory of the posterior negativities and positivities of the event-related brain potential". *International Journal of Psychophysiology*, 105, 111-128.
- [4] Hinton, G. E., McClelland, J. L., & Rumelhart, D. E. (1986). "Distributed representations". In D. E. Rumelhart & J. L. McClelland (Eds.), *Parallel Distributed Processing: Explorations in the Microstructure of Cognition, Volume 1: Foundations* (pp. 77-109). MIT Press.
- [5] Searle, J. R. (1980). "Minds, brains, and programs". *Behavioral and Brain Sciences*, 3(3), 417-424.
- [6] Evans, J. St. B. T. (2003). "In two minds: Dual-process accounts of reasoning". *Trends in Cognitive Sciences*, 7(10), 454-459.
- [7] Friston, K. (2005). "A theory of cortical responses". *Philosophical Transactions of the Royal Society B: Biological Sciences*, 360(1456), 815-836.