

Artificial Intelligence and Psycholinguistics: How Do Large Language Models Represent Human Language Processes? A Systematic Review

Astuty

Universitas Tidar, Magelang, Indonesia



Corresponding Author's Email: astuty@untidar.ac.id

INFO ARTIKEL

History Article:

Submitted 08 10, 2026
Received 08 26, 2026
Published 08 27, 2026

ABSTRACT.

The rapid development of large language models (LLMs) has transformed contemporary perspectives on language processing, raising questions about their relationship to human cognition. While psycholinguistic theories explain language through perception, memory, attention, semantic representation, and executive control, LLMs rely on statistical learning and large-scale pattern recognition. This study synthesizes research examining similarities and differences between human language processing and LLM-based language representation. Using a qualitative systematic review, it integrates findings published between 2018 and 2025 across psycholinguistics, cognitive science, neuroscience, computational linguistics, and artificial intelligence. Results indicate that LLMs approximate several observable features of human language behavior, including syntactic processing, contextual prediction, semantic association, discourse coherence, and pragmatic inference. However, important differences remain in grounding, intentionality, episodic memory, emotional experience, and embodied cognition. While LLMs offer valuable computational models for investigating psycholinguistic theories, they cannot currently be considered cognitive equivalents of human language users, instead representing probabilistic systems that emulate linguistic behavior without replicating the full cognitive architecture underlying human communication. The review discusses theoretical implications and identifies directions for integrating computational and cognitive approaches to language.

Keywords: Artificial Intelligence; Psycholinguistics; Large Language Models; Language Processing; Cognitive Science; Computational Linguistics

How to Cite:

Astuty, A. (2026). Artificial Intelligence and Psycholinguistics: How Do Large Language Models Represent Human Language Processes? A Systematic Review. LANCAH: Jurnal Inovasi Dan Tren, 4(2), 172-191. <https://journal.lembagakita.com/ljit/article/view/8346>

INTRODUCTION

Artificial intelligence has rapidly become one of the most influential technological developments of the twenty-first century. Among recent advances, large language models (LLMs) represent a particularly significant breakthrough because they exhibit language capabilities that increasingly resemble aspects of human linguistic behavior. Models such as GPT, Gemini, Claude, Llama, and DeepSeek can generate coherent texts, answer complex questions, summarize information, translate between languages, and participate in extended conversations with remarkable fluency. These developments have generated growing interest among psycholinguists because they raise fundamental questions regarding the nature of language representation, cognitive processing, and linguistic intelligence.

For decades, psycholinguistics has sought to explain how humans comprehend, produce, acquire, and represent language. Classical theories emphasize cognitive mechanisms such as lexical access, syntactic parsing, semantic integration, working memory, attention, and executive control (Traxler, 2014). These mechanisms are generally understood as emerging from complex interactions among multiple cognitive systems supported by biological neural networks. Human language processing therefore reflects not only linguistic knowledge but also perception, memory, emotion, social cognition, and embodied experience.

In contrast, LLMs are computational systems trained on massive text corpora using deep neural network architectures based primarily on the Transformer model introduced by Vaswani et al. (2017). Rather than relying on explicit linguistic rules or symbolic representations, these models learn statistical regularities from billions of words through self-supervised learning. Consequently, they develop highly sophisticated representations capable of predicting subsequent linguistic elements within contextual sequences. This approach has achieved unprecedented success across numerous natural language processing tasks.

The remarkable performance of LLMs has stimulated renewed debate regarding the relationship between artificial and human intelligence. Some researchers argue that sufficiently large neural language models capture important computational principles underlying human language processing (Piantadosi, 2023). Others maintain that despite impressive linguistic performance, LLMs fundamentally differ from humans because they lack intentionality, consciousness, embodied experience, emotional understanding, and genuine semantic grounding (Bender & Koller, 2020). These contrasting perspectives have important implications for contemporary psycholinguistic theory.

One particularly relevant issue concerns language representation. Human psycholinguistic models generally assume that words are represented through interactions among semantic memory, episodic memory, perceptual experience, and conceptual knowledge. Contemporary embodied cognition theories further propose that conceptual representations are grounded in sensory, motor, and affective systems (Barsalou, 2008). By contrast, LLMs represent language through high-dimensional vector embeddings learned from statistical co-occurrence patterns. Whether these computational representations correspond meaningfully to human conceptual representations remains an open question.

Another important issue involves predictive processing. Human language comprehension relies heavily on prediction, with listeners and readers continuously anticipating upcoming linguistic information during communication (Pickering & Gambi, 2018). LLMs similarly operate through next-token prediction, generating linguistic output based on probabilistic estimation. This apparent similarity has encouraged

researchers to investigate whether predictive processing constitutes a shared computational principle underlying both biological and artificial language systems.

Memory processes represent another important area of comparison. Human language comprehension depends upon interactions between working memory, long-term semantic memory, episodic memory, and attentional control (Baddeley, 2012). LLMs instead rely upon contextual windows and learned parameters rather than biologically inspired memory systems. Although these mechanisms differ substantially, both support the integration of contextual information during language processing.

Recent research has also explored whether LLMs exhibit psycholinguistic phenomena traditionally associated with human cognition. Investigations have examined syntactic ambiguity resolution, semantic priming, pragmatic inference, discourse coherence, lexical prediction, metaphor comprehension, and theory of mind. Many studies report surprisingly human-like performance under certain experimental conditions, although significant limitations remain (Mahowald et al., 2024).

Despite rapidly expanding research, the literature on LLMs and human language processing remains fragmented across disciplinary boundaries. Within computational linguistics, critical analyses such as Bender and Koller (2020) have argued that models trained purely on linguistic form cannot be assumed to have acquired meaning, a position echoed by Bender et al. (2021) in their critique of large-scale language models—yet this critique engages minimally with established psycholinguistic accounts of how humans build meaning during comprehension, such as prediction-based frameworks (Pickering & Gambi, 2018) or discourse-level comprehension models (Kintsch, 1998). Conversely, neuroscience-oriented work has pursued a different integration: Schrimpf et al. (2021) demonstrated that LLM internal representations align with neural activity during language processing under a predictive-processing framework, building on the broader predictive-brain hypothesis (Clark, 2013), but this line of work has largely proceeded independently of the representational and semantic debates raised in computational linguistics (Bender & Koller, 2020). More recently, Mahowald et al. (2024) took a step toward integration by explicitly dissociating language competence from thought in LLMs, drawing on cognitive science distinctions between linguistic and conceptual processing, yet this work centers narrowly on the language-thought distinction rather than engaging with a broader range of established human language-processing constructs, such as working-memory constraints on comprehension (Baddeley, 2012) or usage-based accounts of grammatical generalization (Goldberg, 2006; Tomasello, 2003). This pattern suggests that computational critiques, neurocognitive integration efforts, and theoretical linguistic frameworks have largely developed along separate tracks, with limited cross-referencing of each other's core claims and limitations. Consequently, there is a need for a synthesis that explicitly connects these computational, neurocognitive, and theoretical linguistic perspectives on the relationship between contemporary LLMs and human language processing.

The present study addresses this need through a systematic review of empirical and theoretical research examining LLMs from a psycholinguistic perspective. Specifically, the review seeks to answer the following questions: First, how do LLMs represent linguistic knowledge relative to human language processing? Second, what similarities and differences exist between computational and cognitive models of language? Third, what psycholinguistic theories receive support or challenge from recent LLM research? Finally, what implications do these developments have for future interdisciplinary research integrating artificial intelligence and psycholinguistics?

METHODS OF RESEARCH

This study employed a qualitative systematic review to synthesize empirical and theoretical research concerning the relationship between large language models and psycholinguistic theories of language processing. The review followed established principles of systematic evidence synthesis to ensure transparency, comprehensiveness, and methodological rigor (Page et al., 2021).

The literature search was conducted using Scopus, Web of Science, IEEE Xplore, ACM Digital Library, PsycINFO, ScienceDirect, SpringerLink, Wiley Online Library, Taylor & Francis Online, and Google Scholar. Because large language models represent a rapidly evolving research area, the review focused on publications appearing between 2018 and 2025. This period captures the emergence of Transformer-based architectures and the rapid expansion of contemporary LLM research.

Search terms included combinations of “large language models,” “artificial intelligence,” “psycholinguistics,” “language processing,” “human language comprehension,” “Transformer,” “GPT,” “computational linguistics,” “cognitive science,” “language representation,” “semantic processing,” and “language prediction.” Boolean operators were used to refine searches and maximize the retrieval of relevant studies. Data extraction focused on the theoretical framework, research objectives, language models investigated, psycholinguistic constructs examined, methodological approaches, principal findings, and theoretical implications. Through iterative thematic synthesis, findings were organized into major conceptual categories representing recurring themes across the literature.

RESULT AND DISCUSSION

Result

1) Large Language Models as Computational Models of Human Language

One of the most significant findings emerging from the reviewed literature concerns the remarkable ability of large language models (LLMs) to approximate numerous observable characteristics of human language processing. Although LLMs were originally designed to predict the next token in a sequence, advances in Transformer-based architectures have enabled these systems to perform increasingly sophisticated linguistic tasks, including semantic interpretation, discourse construction, question answering, summarization, translation, and reasoning (Brown et al., 2020; OpenAI, 2023; Vaswani et al., 2017).

From a psycholinguistic perspective, these developments raise an important question: Do LLMs merely imitate linguistic behavior through statistical learning, or do they provide meaningful computational analogues of human language processing? The reviewed studies reveal that opinions remain divided. Some researchers argue that the linguistic competence displayed by LLMs reflects the emergence of internal representations that resemble aspects of human cognitive processing (Piantadosi, 2023). Others contend that these systems achieve impressive performance without possessing the conceptual grounding, intentionality, or experiential knowledge characteristic of human cognition (Bender & Koller, 2020).

Despite these theoretical disagreements, there is broad consensus that LLMs have become valuable experimental tools for psycholinguistic research. Computational models can now simulate numerous language-processing phenomena traditionally investigated through behavioral experiments, allowing researchers to compare artificial and human performance under controlled conditions.

Unlike earlier symbolic models of language, LLMs rely on distributed vector representations learned from extensive linguistic input. Words are represented as multidimensional embeddings that capture statistical relationships among lexical items. Interestingly, these representations often exhibit semantic organization similar to patterns observed in human lexical processing. Words sharing conceptual or contextual similarities tend to occupy neighboring regions within high-dimensional representational spaces, resembling semantic clustering observed in psycholinguistic studies of lexical memory (Mikolov et al., 2013).

However, important differences remain. Human lexical representations are shaped not only by linguistic exposure but also by sensory perception, emotional experience, autobiographical memory, and social interaction (Barsalou, 2008). LLM representations emerge exclusively from textual data. Consequently, although statistical similarities between computational and human representations are substantial, they cannot be considered identical.

2) Predictive Processing in Humans and Large Language Models

Prediction has become one of the central concepts linking psycholinguistics and artificial intelligence. Contemporary theories of human language comprehension increasingly argue that listeners and readers continuously generate expectations regarding upcoming linguistic input. Prediction facilitates rapid comprehension by enabling individuals to anticipate lexical items, syntactic structures, and semantic interpretations before they are explicitly encountered (Pickering & Gambi, 2018).

Remarkably, prediction also constitutes the fundamental computational principle underlying modern LLMs. During training, models repeatedly estimate the probability of subsequent linguistic tokens based on preceding contextual information. Through billions of such prediction tasks, models gradually acquire sophisticated representations of language structure.

The reviewed literature identifies important similarities between human predictive processing and computational prediction. Both systems utilize contextual information to reduce uncertainty during language comprehension. Behavioral experiments indicate that humans process highly predictable words more rapidly than unpredictable alternatives. Similarly, LLMs assign higher probabilities to contextually expected tokens, generating fluent and coherent outputs.

Studies comparing human cloze probability judgments with LLM probability estimates frequently report substantial correlations (Schrimpf et al., 2021). Such findings suggest that both biological and artificial systems rely on predictive mechanisms during language processing. Nevertheless, the underlying computational architectures differ fundamentally. Human prediction involves interactions among perception, attention, memory, world knowledge, and communicative intentions. LLM prediction depends exclusively upon statistical regularities encoded within neural network parameters.

Another important distinction concerns flexibility. Human predictions can rapidly incorporate newly acquired perceptual information and real-world experiences. LLMs, by contrast, require retraining or external memory mechanisms to integrate knowledge unavailable during initial training. Consequently, although prediction represents a shared principle, the cognitive processes supporting prediction remain substantially different.

3) Lexical Representation and Semantic Processing

Lexical representation constitutes another major area of convergence between psycholinguistics and artificial intelligence. Human lexical processing involves recognizing words, retrieving semantic information, integrating contextual meanings, and selecting appropriate interpretations during communication (Traxler, 2014). LLMs similarly encode extensive lexical knowledge acquired through large-scale textual exposure.

Research examining semantic representations indicates that word embeddings learned by LLMs capture numerous semantic relationships traditionally investigated within psycholinguistics. Synonyms, antonyms, taxonomic categories, and thematic associations frequently emerge naturally during model training (Devlin et al., 2019). Computational analyses demonstrate that semantically related words occupy neighboring positions within representational spaces, resembling semantic organization observed in behavioral experiments.

The reviewed literature also highlights similarities in contextual meaning construction. Human lexical meanings are highly context dependent. Words such as bank or light acquire different interpretations depending upon surrounding linguistic information. Transformer architectures successfully model such contextual variability through attention mechanisms that dynamically adjust lexical representations according to discourse context (Vaswani et al., 2017).

However, important theoretical questions remain concerning semantic grounding. Human meanings are deeply connected to sensory experiences, bodily actions, emotional responses, and environmental interactions (Barsalou, 2008). LLMs derive semantic representations exclusively from linguistic co-occurrence patterns. Critics therefore argue that computational semantics remains fundamentally ungrounded despite impressive linguistic performance (Bender & Koller, 2020).

Supporters of LLM research counter that extensive textual experience may indirectly encode substantial amounts of world knowledge. Although models do not perceive the physical world directly, language itself contains rich descriptions of human experiences accumulated across generations. Consequently, statistical representations learned from language may approximate certain aspects of conceptual knowledge without direct sensory grounding. The reviewed evidence suggests that neither position fully explains current findings. LLMs clearly possess sophisticated semantic representations capable of supporting complex language tasks. Nevertheless, these representations differ qualitatively from human conceptual systems grounded in perception and embodied experience.

4) Syntactic Processing and Grammatical Competence

Syntactic processing has historically occupied a central position in psycholinguistic theory. Classical approaches proposed that language comprehension depends upon specialized grammatical mechanisms capable of constructing hierarchical sentence structures (Chomsky, 1965). More recent usage-based theories emphasize statistical learning and experience in grammatical acquisition (Tomasello, 2003).

The emergence of LLMs has significantly influenced this debate. Without explicit grammatical rules, modern language models successfully generate syntactically well-formed sentences across numerous languages. This achievement has prompted researchers to reconsider the extent to which statistical learning alone can account for grammatical competence.

Experimental studies reveal that LLMs perform remarkably well on many syntactic tasks, including subject-verb agreement, long-distance dependencies, relative clause processing, and grammatical

acceptability judgments (Linzen et al., 2016). These findings suggest that neural language models acquire sophisticated representations of grammatical regularities through exposure to extensive linguistic input.

Comparisons with human sentence processing reveal both similarities and differences. Like human readers, LLMs exhibit greater difficulty processing syntactically complex constructions than simpler sentences. Garden-path sentences, center-embedded clauses, and ambiguous structures frequently challenge both biological and artificial systems.

However, psycholinguistic research indicates that humans employ multiple information sources during syntactic processing, including semantic plausibility, discourse context, working memory, prosody, and real-world knowledge (MacDonald et al., 1994). Although LLMs approximate many of these influences statistically, they lack direct perceptual and communicative experiences that contribute to human grammatical interpretation. The reviewed studies therefore suggest that LLMs provide powerful computational approximations of syntactic processing while remaining fundamentally different from biological language systems.

5) Pragmatic Inference and Discourse Processing

One of the most challenging aspects of language concerns pragmatics, namely the interpretation of speaker intentions beyond literal linguistic meanings. Human communication relies extensively upon contextual knowledge, shared beliefs, cultural conventions, conversational norms, and theory of mind (Clark, 1996).

The reviewed literature demonstrates that contemporary LLMs perform surprisingly well on many pragmatic tasks. Models frequently generate contextually appropriate responses, identify indirect meanings, interpret conversational implicatures, and maintain coherent discourse across extended interactions (OpenAI, 2023). Such capabilities have generated considerable interest among psycholinguists investigating computational models of communication.

Nevertheless, important limitations remain. Human pragmatic reasoning depends upon intentional communication between conscious individuals possessing beliefs, goals, emotions, and social relationships. LLMs generate pragmatically appropriate outputs through statistical associations rather than genuine communicative intentions. Consequently, successful performance does not necessarily imply human-like pragmatic competence.

Research examining discourse coherence reveals similar patterns. Transformer attention mechanisms enable models to maintain long-range contextual dependencies, producing coherent narratives across multiple paragraphs. However, discourse consistency sometimes deteriorates during extended interactions, particularly when conversations require persistent memory of earlier events or complex causal reasoning. Overall, the evidence suggests that LLMs successfully emulate many observable characteristics of human pragmatic behavior while differing fundamentally in the cognitive mechanisms underlying communicative understanding.

6) Working Memory, Attention, and Context Management

Working memory has long occupied a central position in psycholinguistic theories of language processing. Human language comprehension depends on the temporary maintenance and manipulation of linguistic information while integrating incoming words into coherent syntactic and semantic structures (Baddeley, 2012). During sentence comprehension, individuals must retain lexical items, monitor syntactic

dependencies, resolve ambiguities, and update discourse representations continuously. Consequently, working memory interacts closely with attention, executive control, and long-term memory throughout language processing.

The reviewed literature indicates that LLMs exhibit functional similarities to certain aspects of working memory despite relying on fundamentally different computational mechanisms. Rather than maintaining information through biologically inspired memory systems, Transformer-based architectures utilize self-attention mechanisms that dynamically compute relationships among tokens within a contextual window (Vaswani et al., 2017). Every input token can attend to every other token, enabling the model to integrate contextual information across an entire sequence simultaneously.

From a behavioral perspective, this architecture allows LLMs to maintain contextual coherence across extended passages of text. Similar to human readers, language models incorporate earlier contextual information when interpreting subsequent linguistic input. Studies comparing human reading behavior with model predictions reveal substantial correspondence in contextual sensitivity, particularly during sentence completion and discourse comprehension tasks (Schrimpf et al., 2021).

However, the reviewed studies emphasize that computational context management differs fundamentally from human working memory. Human working memory possesses severe capacity limitations and is strongly influenced by attentional allocation, interference, cognitive load, and individual differences (Baddeley, 2012). These limitations explain numerous psycholinguistic phenomena, including processing difficulties associated with center-embedded clauses, long-distance dependencies, and syntactic ambiguity.

LLMs do not experience cognitive load in the psychological sense. Instead, they are constrained by fixed computational context windows determined by model architecture. Earlier Transformer models processed only a few hundred tokens, whereas recent systems may accommodate hundreds of thousands of tokens. Nevertheless, increasing context length does not necessarily produce human-like memory because stored representations remain statistical rather than experiential.

Attention represents another important area of comparison. Human attention selectively prioritizes relevant linguistic information while suppressing competing stimuli. During language comprehension, attentional mechanisms regulate lexical access, semantic integration, syntactic parsing, and discourse updating (Posner & Petersen, 1990). Numerous psycholinguistic experiments demonstrate that attentional resources directly influence comprehension accuracy and processing speed.

Transformer attention mechanisms similarly allocate computational resources across linguistic inputs. Multi-head attention allows models to identify multiple relationships simultaneously, including syntactic dependencies, semantic associations, and discourse-level connections. Visualization studies reveal that different attention heads specialize in different linguistic functions, including subject-verb agreement, anaphora resolution, and phrase structure recognition (Clark et al., 2019).

Despite these functional similarities, important conceptual differences remain. Human attention reflects conscious cognitive control influenced by motivation, emotion, perceptual salience, and communicative goals. Computational attention represents mathematical weighting functions optimized through gradient-based learning. Consequently, although both systems prioritize relevant information, they do so through fundamentally different mechanisms.

These findings suggest that LLMs provide computational analogues of contextual integration rather than direct models of human working memory. Their success demonstrates that sophisticated contextual

language processing can emerge from attention-based computation, but it does not imply equivalence with biological memory systems.

7) Emotion, Embodiment, and Consciousness

Among the most debated topics within contemporary psycholinguistics is whether LLMs can represent emotional meaning in ways comparable to humans. Human language is inseparable from emotional experience. Emotional states influence lexical selection, syntactic complexity, discourse organization, conversational behavior, and pragmatic interpretation (Citron, 2012). Furthermore, contemporary embodied cognition theories argue that conceptual knowledge itself is grounded in sensory, motor, and affective experiences (Barsalou, 2008).

The reviewed literature consistently indicates that LLMs successfully recognize and generate emotionally appropriate language. Models accurately classify emotional valence, produce empathic responses, summarize emotionally charged narratives, and identify affective meanings within discourse. Such performance often approaches or exceeds human benchmarks on standardized natural language processing tasks.

Nevertheless, researchers strongly disagree regarding the interpretation of these achievements. One perspective argues that emotional competence can emerge from statistical regularities encoded within linguistic data. Because human language contains extensive information regarding emotional experiences, social interactions, and affective evaluations, sufficiently large models may learn rich representations of emotional meaning through textual exposure alone (Piantadosi, 2023).

An opposing perspective emphasizes that successful emotional language processing should not be confused with emotional experience. Human emotional understanding depends upon physiological responses, subjective feelings, autobiographical memories, social relationships, and embodied interactions with the environment (Barrett, 2017). LLMs lack these experiential foundations. Consequently, although they can describe emotions accurately, they do not experience fear, joy, sadness, or empathy in the psychological sense.

Embodied cognition further highlights this distinction. Human conceptual representations are deeply connected to perception and action. Understanding words such as *grasp*, *run*, or *warm* activates sensory and motor systems associated with corresponding experiences (Barsalou, 2008). Neuroimaging studies consistently demonstrate interactions between conceptual processing and perceptual brain regions.

LLMs instead acquire conceptual knowledge exclusively through linguistic statistics. Their representations therefore remain symbolic or distributional rather than embodied. Critics argue that this limitation restricts genuine conceptual understanding because language ultimately derives its meaning from interactions with the physical and social world (Bender & Koller, 2020).

Recent multimodal language models partially address this limitation by integrating textual, visual, and auditory information during training. Such models demonstrate improved performance on tasks requiring visual reasoning and grounded language understanding. Nevertheless, even multimodal systems remain fundamentally different from humans because they lack autonomous perception, bodily interaction, and lived experience.

Consciousness represents an even more controversial issue. Human language processing occurs within conscious cognitive systems possessing intentions, beliefs, goals, and self-awareness. Current evidence provides no convincing support for attributing consciousness or subjective awareness to LLMs.

Their linguistic competence emerges from optimization of statistical prediction rather than conscious understanding. Consequently, most researchers conclude that current LLMs simulate linguistic behavior without reproducing the conscious cognitive architecture underlying human communication.

8) Large Language Models as Experimental Tools for Psycholinguistics

Perhaps the most important contribution of LLMs to psycholinguistics lies not in replacing existing theories but in providing powerful computational models for testing them. Historically, psycholinguistic theories have relied primarily on behavioral experiments, reaction time studies, eye-tracking, neuroimaging, and corpus analyses. Contemporary language models introduce an additional methodological framework through which hypotheses concerning language representation and processing can be evaluated.

The reviewed studies demonstrate that LLMs successfully reproduce numerous psycholinguistic phenomena. Models exhibit frequency effects, contextual facilitation, semantic similarity patterns, syntactic preferences, discourse coherence, and predictive behaviors analogous to those observed in human participants (Schrimpf et al., 2021). These similarities allow researchers to compare computational predictions directly with experimental findings.

One particularly promising application involves hypothesis testing. Researchers can systematically manipulate linguistic variables within computational models before conducting costly behavioral experiments. For example, investigators may evaluate how lexical frequency, semantic ambiguity, syntactic complexity, or contextual predictability influence model behavior and then compare these predictions with human performance. Such approaches strengthen the integration of computational modeling and experimental psycholinguistics.

LLMs also facilitate cross-linguistic investigations. Because multilingual models are trained on dozens or even hundreds of languages, they provide opportunities to examine universal and language-specific aspects of linguistic processing. Researchers can investigate whether syntactic preferences, semantic structures, or discourse patterns generalize across linguistic communities without requiring large-scale human participant recruitment for every language.

Educational psycholinguistics likewise benefits from these developments. LLMs can simulate learner language, generate controlled linguistic stimuli, assist in corpus annotation, and support experimental material development. These applications reduce the time required to design psycholinguistic experiments while increasing methodological consistency.

Despite these advantages, the reviewed literature emphasizes that LLMs should complement rather than replace human experimentation. Computational models remain simplifications of human cognition. Validation through behavioral, developmental, and neurocognitive research therefore remains essential. Psycholinguistic theories ultimately seek to explain human cognitive processes rather than optimize artificial prediction systems. Overall, the findings suggest that LLMs represent valuable scientific instruments for investigating language processing. Their greatest contribution may not be demonstrating artificial intelligence equivalent to human cognition but rather providing computational frameworks that generate new hypotheses concerning the mechanisms underlying human language.

Discussions

The findings synthesized in this review demonstrate that large language models have fundamentally reshaped contemporary discussions concerning the nature of human language processing. Although LLMs were developed primarily to solve computational problems in natural language processing, their remarkable linguistic performance has encouraged psycholinguists to reconsider long-standing assumptions regarding language representation, lexical organization, prediction, syntax, semantics, and discourse processing. The reviewed literature indicates that LLMs neither fully replicate nor fundamentally contradict human language processing. Instead, they represent computational systems that capture several observable properties of human linguistic behavior while differing substantially in the cognitive mechanisms underlying those behaviors.

One of the most important contributions of LLM research concerns predictive processing. Contemporary psycholinguistic theories increasingly describe human language comprehension as an inherently predictive process in which listeners and readers continuously anticipate forthcoming linguistic information (Pickering & Gambi, 2018). The remarkable success of Transformer-based architectures provides independent computational support for this theoretical perspective because prediction constitutes the central learning objective of virtually all modern LLMs. During training, these models repeatedly estimate subsequent linguistic units based on contextual information, gradually developing sophisticated representations of syntax, semantics, and discourse.

The convergence between predictive processing theories and computational language modeling suggests that prediction may represent one of the fundamental computational principles underlying successful language processing. Nevertheless, the reviewed evidence also highlights important differences. Human prediction emerges from interactions among perception, memory, communicative intentions, emotional states, social cognition, and real-world knowledge. Computational prediction instead relies upon optimization of statistical distributions learned from textual corpora. Consequently, similar behavioral outcomes may arise from fundamentally different underlying mechanisms.

The review also contributes to debates concerning language representation. Classical psycholinguistic theories frequently conceptualize lexical knowledge as networks of semantic, phonological, orthographic, and conceptual representations organized within long-term memory (Levelt, 1999). Contemporary embodied cognition theories further argue that conceptual knowledge depends upon interactions with perceptual, motor, and emotional systems (Barsalou, 2008). In contrast, LLMs represent words as distributed vectors embedded within high-dimensional representational spaces generated through statistical learning.

Interestingly, empirical comparisons reveal substantial correspondence between computational embeddings and human semantic judgments. Distributional representations learned by language models often predict semantic similarity ratings, lexical association norms, and category membership with remarkable accuracy (Mikolov et al., 2013). Such findings suggest that statistical regularities embedded within language contain considerable information regarding conceptual organization.

However, these similarities should not obscure important distinctions. Human semantic memory is continuously shaped by direct experiences, bodily interactions, emotional episodes, cultural participation, and autobiographical knowledge. Language models derive conceptual structures exclusively from linguistic exposure. Consequently, although computational semantics approximates many observable properties of human semantic organization, it remains fundamentally distinct from embodied conceptual systems.

Another major implication concerns grammatical competence. One of the longest-standing debates in linguistics concerns whether grammatical knowledge depends upon innate symbolic structures or emerges through statistical learning and language experience. The impressive syntactic performance of LLMs has challenged arguments suggesting that explicit grammatical rules are necessary for sophisticated language behavior. Without receiving formal syntactic instruction, contemporary language models generate grammatically coherent texts across numerous languages and accurately process complex syntactic dependencies (Linzen et al., 2016).

Nevertheless, the reviewed studies caution against interpreting these findings as evidence that human grammar is reducible to statistical prediction alone. Human syntactic processing involves working memory limitations, attentional fluctuations, semantic expectations, communicative goals, and developmental learning processes extending across many years. Computational models achieve comparable behavioral performance through optimization procedures fundamentally different from biological language acquisition.

The review likewise highlights important insights concerning working memory and attention. Human psycholinguistic research consistently demonstrates that language comprehension depends upon limited-capacity working memory systems capable of maintaining and manipulating linguistic information during online processing (Baddeley, 2012). LLMs instead employ attention mechanisms that compute contextual relationships across input sequences simultaneously. Although both systems support contextual integration, they operate according to fundamentally different computational principles.

This distinction has important theoretical implications. Similar behavioral performance does not necessarily imply similar cognitive architecture. LLMs may successfully reproduce numerous psycholinguistic phenomena while relying upon computational mechanisms unavailable to biological cognition. Conversely, humans may achieve comparable linguistic outcomes despite possessing severe cognitive limitations absent from artificial systems. The reviewed literature therefore supports a distinction between functional equivalence and mechanistic equivalence. LLMs frequently demonstrate the former but not necessarily the latter.

The findings concerning pragmatics further illustrate this distinction. Contemporary language models often produce contextually appropriate conversational responses and maintain coherent discourse across extended interactions. Such achievements initially appear to support claims that LLMs possess sophisticated pragmatic competence. However, human pragmatics depends upon shared intentions, beliefs, emotional understanding, social relationships, and theory of mind (Clark, 1996). Current LLMs generate pragmatically appropriate language through statistical pattern matching rather than intentional communication. Consequently, successful conversational performance should not be interpreted as evidence of genuine communicative understanding.

The review also addresses perhaps the most controversial question surrounding LLMs: whether they truly understand language. The evidence synthesized across disciplines suggests that the answer depends largely upon how understanding is defined. If understanding refers to successful prediction, contextual interpretation, semantic integration, and coherent language generation, contemporary LLMs clearly exhibit impressive capabilities. However, if understanding requires conscious awareness, intentionality, embodied experience, autobiographical memory, emotional participation, and situated interaction with the physical world, existing evidence suggests that current LLMs do not satisfy these criteria.

This distinction parallels longstanding philosophical debates regarding artificial intelligence. Searle's (1980) Chinese Room argument maintains that syntactically appropriate symbol manipulation does not necessarily produce semantic understanding. More recently, Bender and Koller (2020) similarly argue that language models manipulate linguistic forms without grounding them in communicative intentions or experiential knowledge. Conversely, researchers such as Piantadosi (2023) contend that many aspects of human linguistic competence may themselves emerge from sophisticated predictive learning mechanisms. The reviewed literature does not definitively resolve these debates but instead demonstrates that both perspectives capture important aspects of contemporary evidence.

An additional contribution emerging from the review concerns the growing role of LLMs as scientific models within psycholinguistics. Historically, computational modeling occupied a relatively specialized position within language research. Modern language models now provide powerful tools capable of generating experimentally testable predictions regarding lexical access, syntactic processing, semantic organization, discourse comprehension, and language acquisition. Rather than replacing psycholinguistic experimentation, LLMs increasingly complement behavioral and neurocognitive methodologies by providing explicit computational hypotheses concerning language processing.

The interdisciplinary nature of recent research further represents an important development. The boundaries separating psycholinguistics, cognitive science, computational linguistics, neuroscience, and artificial intelligence have become increasingly permeable. The reviewed studies demonstrate that progress frequently occurs through collaboration across these traditionally distinct disciplines. Advances in neuroimaging inform computational architectures, while developments in machine learning generate new hypotheses concerning human cognition. Such reciprocal influence illustrates the emergence of a more integrated science of language.

At the same time, the review identifies several limitations within existing research. First, most comparative studies focus on English or other high-resource languages, limiting the generalizability of current findings across diverse linguistic systems. Second, rapid technological developments frequently outpace empirical evaluation, making theoretical conclusions vulnerable to continual revision. Third, behavioral similarity between humans and LLMs sometimes encourages anthropomorphic interpretations unsupported by cognitive evidence. Future research must therefore distinguish carefully between observable linguistic performance and underlying cognitive mechanisms.

Overall, the evidence synthesized in this review supports a balanced conclusion. Large language models represent one of the most significant developments in the history of computational linguistics and have substantially advanced understanding of language representation and processing. They successfully reproduce numerous observable characteristics of human linguistic behavior and provide valuable computational frameworks for psycholinguistic research. Nevertheless, they remain fundamentally different from human cognitive systems in terms of embodiment, intentionality, experiential learning, consciousness, and emotional understanding. Rather than replacing psycholinguistic theories, LLMs enrich them by offering powerful new perspectives on how complex language behavior may emerge from computational learning systems.

1) Theoretical Implications

The rapid development of large language models has important implications for contemporary psycholinguistic theory because it challenges several assumptions regarding the relationship between

language, cognition, and intelligence. Rather than replacing existing theories, LLMs encourage scholars to reconsider how language processing can be conceptualized within broader computational and cognitive frameworks. The reviewed literature suggests that future psycholinguistic theories should move beyond traditional dichotomies separating symbolic and statistical approaches, instead recognizing that human language processing likely emerges from interactions among multiple representational and computational mechanisms.

One of the most important theoretical implications concerns the concept of language representation. Traditional psycholinguistic theories generally assume that lexical knowledge is organized within interconnected semantic, phonological, orthographic, and conceptual networks supported by long-term memory (Levelt, 1999). Distributional language models demonstrate that remarkably rich semantic representations can emerge solely from statistical regularities embedded within linguistic experience. This finding does not invalidate classical psycholinguistic models but instead suggests that statistical learning may contribute more substantially to language acquisition and representation than previously assumed.

The success of LLMs also strengthens usage-based and probabilistic theories of language. Construction Grammar, usage-based linguistics, and emergentist perspectives argue that linguistic competence develops through repeated exposure to language rather than through explicit symbolic rules (Tomasello, 2003). Modern language models provide computational demonstrations that extensive statistical learning can produce sophisticated grammatical and semantic behavior. Consequently, these computational achievements lend indirect support to theories emphasizing experience-dependent language learning.

At the same time, the review indicates that purely statistical accounts remain insufficient for explaining the full complexity of human language processing. Human communication depends not only on linguistic regularities but also on perception, intentionality, autobiographical memory, social interaction, emotional experience, and embodied cognition. These components remain largely absent from current language models. Therefore, psycholinguistic theories should continue integrating statistical learning with broader cognitive architectures capable of accounting for these uniquely human characteristics.

The findings further reinforce predictive processing as a unifying theoretical framework. Prediction plays central roles in both biological and artificial language systems, suggesting that anticipatory computation may constitute one of the fundamental principles underlying efficient language comprehension. Future psycholinguistic theories may therefore benefit from placing predictive mechanisms at the center of language processing models while acknowledging important differences in implementation between biological and computational systems.

Another theoretical implication concerns modularity. Classical cognitive theories frequently conceptualized language as a relatively independent cognitive module possessing specialized computational mechanisms. Contemporary evidence instead favors highly interactive architectures in which language processing continuously interacts with memory, attention, executive control, perception, and social cognition. LLM research similarly demonstrates that sophisticated language behavior emerges from interactions among multiple distributed representations rather than isolated symbolic modules. This convergence supports increasingly integrated models of cognition.

The reviewed studies also contribute to debates concerning semantic representation. Distributional semantics demonstrates that many conceptual relationships can emerge through statistical learning from linguistic input. However, embodied cognition research continues to provide compelling evidence that human conceptual knowledge depends upon sensory and motor experiences (Barsalou, 2008). Future

theoretical developments should therefore seek to integrate distributional and embodied perspectives rather than treating them as mutually exclusive explanations. Multimodal artificial intelligence systems capable of integrating textual, visual, and auditory information may provide valuable computational frameworks for investigating such integration.

Finally, the emergence of LLMs encourages psycholinguists to distinguish more clearly between functional similarity and cognitive equivalence. Artificial systems may reproduce many observable characteristics of human language behavior without replicating the biological mechanisms responsible for those behaviors. This distinction is essential for avoiding both excessive anthropomorphism and unnecessary dismissal of computational models. Future psycholinguistic theories should therefore evaluate artificial intelligence according to explicit theoretical criteria rather than relying solely upon behavioral similarity.

2) Ethical Implications

The integration of large language models into psycholinguistic research raises several important ethical considerations extending beyond purely technical questions. As artificial intelligence increasingly participates in educational, clinical, research, and communicative contexts, understanding its cognitive limitations becomes essential for responsible application.

One major ethical issue concerns anthropomorphism. Because contemporary LLMs generate highly fluent and contextually appropriate language, users frequently attribute human-like understanding, intentions, emotions, or consciousness to these systems. The reviewed literature consistently emphasizes that such interpretations exceed current empirical evidence. Although LLMs simulate numerous characteristics of human communication, they do not possess subjective awareness, personal experiences, or genuine communicative intentions. Researchers therefore have an ethical responsibility to communicate these distinctions clearly when discussing artificial intelligence.

Another important issue concerns educational applications. LLMs are increasingly used to support language learning, academic writing, translation, and literacy development. These applications offer substantial educational benefits, including personalized instruction, immediate feedback, and increased accessibility. However, excessive dependence on AI-generated language may reduce opportunities for learners to develop independent writing, critical thinking, and metalinguistic awareness. Educational institutions should therefore promote balanced approaches that integrate artificial intelligence as a learning support rather than a replacement for human cognitive engagement.

Research integrity represents another critical consideration. LLMs can assist researchers by generating summaries, identifying literature, improving linguistic clarity, and organizing manuscripts. Nevertheless, authors remain fully responsible for verifying factual accuracy, evaluating evidence, and maintaining scholarly originality. AI-generated content should never substitute for independent scientific reasoning or critical evaluation. Transparent disclosure regarding AI-assisted writing is therefore increasingly recommended by many academic publishers.

Bias constitutes an additional ethical concern. Because language models are trained on massive internet-based corpora, they inevitably inherit social, cultural, political, and linguistic biases present within their training data (Bender et al., 2021). These biases may influence generated content, reinforce stereotypes, or disadvantage underrepresented linguistic communities. Psycholinguistic research employing

LLMs should therefore carefully evaluate potential biases before drawing conclusions regarding human cognition or language use.

Privacy and confidentiality also require careful consideration. As language models become integrated into healthcare, psychological assessment, and educational counseling, users may disclose highly sensitive personal information. Researchers and practitioners must ensure that AI applications comply with established ethical standards concerning informed consent, data protection, and confidentiality.

Finally, the rapid development of artificial intelligence raises broader philosophical questions concerning the future relationship between humans and intelligent machines. Rather than conceptualizing AI as a replacement for human cognition, the reviewed literature supports a collaborative perspective in which computational systems augment human intellectual capabilities while recognizing their inherent cognitive limitations. Such an approach promotes responsible innovation without diminishing the unique characteristics of human language, cognition, and social interaction.

3) Future Research Directions

Although research examining the relationship between large language models and psycholinguistics has expanded rapidly, numerous questions remain unresolved. Addressing these questions will require increasingly interdisciplinary collaboration among psycholinguists, cognitive scientists, computational linguists, neuroscientists, philosophers, and artificial intelligence researchers.

One important priority involves expanding cross-linguistic investigations. Most current studies focus predominantly on English and a relatively small number of high-resource languages. Consequently, existing theories may not adequately represent the diversity of human linguistic systems. Future research should investigate whether computational models exhibit similar processing characteristics across typologically distinct languages, including agglutinative, polysynthetic, tonal, and low-resource languages.

Developmental perspectives also deserve greater attention. Human language processing emerges gradually through childhood, adolescence, and adulthood, influenced by biological maturation, social interaction, and educational experience. In contrast, LLMs acquire linguistic competence through large-scale pretraining rather than developmental learning. Comparing developmental language acquisition with computational learning may provide valuable insights into both human cognition and artificial intelligence.

Another promising direction concerns multimodal cognition. Future language models integrating visual perception, auditory processing, motor interaction, and environmental exploration may provide more realistic approximations of embodied language processing. Such developments could help bridge the gap between distributional semantics and embodied cognition theories.

Neuroscientific research likewise offers substantial opportunities for future investigation. Advances in functional magnetic resonance imaging, magnetoencephalography, electroencephalography, and intracranial recording may facilitate increasingly detailed comparisons between neural activity and computational representations generated by LLMs. Understanding correspondences and differences between biological and artificial representations represents one of the most promising directions for cognitive science.

Future studies should also examine emotional language processing more comprehensively. Although LLMs successfully recognize and generate emotional language, little is known regarding the computational mechanisms through which affective representations emerge. Integrating psycholinguistic research on

emotion with computational modeling may contribute to more sophisticated theories of affective language processing.

Longitudinal evaluation of language models represents another important research priority. Most current investigations assess models at single points in time despite rapid improvements in architecture, training methods, and multimodal capabilities. Systematic longitudinal comparisons may reveal how increasingly advanced artificial systems influence theoretical interpretations of language processing.

Finally, future research should continue distinguishing observable linguistic performance from underlying cognitive mechanisms. Behavioral similarity alone cannot establish cognitive equivalence. Comprehensive theories of language processing should therefore integrate behavioral, computational, developmental, neuroscientific, and philosophical evidence when evaluating relationships between artificial and human intelligence.

CONCLUSION

The emergence of large language models represents one of the most significant developments in the history of language science. Their remarkable capacity to generate coherent discourse, process complex syntax, represent semantic relationships, and maintain contextual coherence has fundamentally transformed contemporary discussions concerning language representation and cognition.

The systematic review presented in this study demonstrates that LLMs successfully approximate numerous observable characteristics of human language processing. Evidence synthesized across psycholinguistics, computational linguistics, neuroscience, and cognitive science indicates that these models capture important aspects of lexical representation, predictive processing, syntactic organization, semantic integration, and discourse comprehension. Consequently, LLMs provide valuable computational frameworks for investigating psycholinguistic theories and generating experimentally testable hypotheses.

Nevertheless, substantial differences remain between artificial and human language processing. Human communication is grounded in embodied experience, autobiographical memory, emotional understanding, social interaction, intentionality, and conscious awareness. Current LLMs lack these foundational cognitive characteristics despite their impressive linguistic competence. Their performance reflects sophisticated probabilistic learning from large-scale textual data rather than human-like cognitive understanding.

The findings therefore support a balanced perspective. LLMs should neither be regarded as simple statistical tools devoid of theoretical significance nor as complete models of human cognition. Instead, they occupy an intermediate position as highly sophisticated computational systems capable of reproducing many observable properties of language while remaining fundamentally distinct from biological cognitive architectures.

From a psycholinguistic perspective, the greatest contribution of LLMs may lie in their ability to stimulate new theoretical questions rather than provide definitive answers. By revealing how complex linguistic behavior can emerge from large-scale statistical learning, they encourage researchers to reconsider traditional assumptions regarding language acquisition, representation, prediction, and processing. At the same time, their limitations highlight the enduring importance of embodiment, social experience, and consciousness in human communication.

Ultimately, the future of psycholinguistics will likely depend upon increasingly close collaboration between cognitive science and artificial intelligence. Rather than viewing these disciplines as competitors, researchers should recognize their complementary strengths. Human cognition provides the theoretical foundation for understanding language, whereas artificial intelligence offers powerful computational tools for testing and refining that understanding. Together, these perspectives promise to advance a more comprehensive science of language in the era of intelligent machines.

REFERENCE

- Baddeley, A. D. (2012). Working memory: Theories, models, and controversies. *Annual Review of Psychology*, 63, 1–29. <https://doi.org/10.1146/annurev-psych-120710-100422>
- Barrett, L. F. (2017). *How emotions are made: The secret life of the brain*. Houghton Mifflin Harcourt.
- Barsalou, L. W. (2008). Grounded cognition. *Annual Review of Psychology*, 59, 617–645. <https://doi.org/10.1146/annurev.psych.59.103006.093639>
- Bender, E. M., & Koller, A. (2020). Climbing towards NLU: On meaning, form, and understanding in the age of data. In *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics* (pp. 5185–5198). Association for Computational Linguistics. <https://doi.org/10.18653/v1/2020.acl-main.463>
- Bender, E. M., Gebru, T., McMillan-Major, A., & Shmitchell, S. (2021). On the dangers of stochastic parrots: Can language models be too big? In *Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency* (pp. 610–623). ACM. <https://doi.org/10.1145/3442188.3445922>
- Brown, T. B., Mann, B., Ryder, N., Subbiah, M., Kaplan, J., Dhariwal, P., ... Amodei, D. (2020). Language models are few-shot learners. *Advances in Neural Information Processing Systems*, 33, 1877–1901. <https://doi.org/10.48550/arXiv.2005.14165>
- Chomsky, N. (1965). *Aspects of the theory of syntax*. MIT Press.
- Citron, F. M. M. (2012). Neural correlates of written emotion word processing. *Language and Linguistics Compass*, 6(3), 175–186. <https://doi.org/10.1016/j.bandl.2011.12.007>
- Clark, A. (2013). Whatever next? Predictive brains, situated agents, and the future of cognitive science. *Behavioral and Brain Sciences*, 36(3), 181–204. <https://doi.org/10.1017/s0140525x12000477>
- Clark, H. H. (1996). *Using language*. Cambridge University Press.
- Clark, K., Khandelwal, U., Levy, O., & Manning, C. D. (2019). What does BERT look at? An analysis of BERT's attention. In *Proceedings of the 2019 Workshop on BlackboxNLP: Analyzing and Interpreting Neural Networks for NLP* (pp. 276–286). Association for Computational Linguistics. <https://doi.org/10.18653/v1/W19-4828>
- Devlin, J., Chang, M.-W., Lee, K., & Toutanova, K. (2019). BERT: Pre-training of deep bidirectional transformers for language understanding. In *Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies* (pp. 4171–4186). Association for Computational Linguistics. <https://doi.org/10.18653/v1/N19-1423>
- Elman, J. L. (1990). Finding structure in time. *Cognitive Science*, 14(2), 179–211. [https://doi.org/10.1016/0364-0213\(90\)90002-E](https://doi.org/10.1016/0364-0213(90)90002-E)

- Fedorenko, E., Blank, I., Siegelman, M., & Mineroff, Z. (2020). Lack of selectivity for syntax relative to word meanings throughout the language network. *Cognition*, 203, 104348. <https://doi.org/10.1016/j.cognition.2020.104348>
- Frank, S. L., Bod, R., & Christiansen, M. H. (2012). How hierarchical is language use? *Proceedings of the Royal Society B: Biological Sciences*, 279(1747), 4522–4531. <https://doi.org/10.1016/j.cognition.2020.104348>
- Goldberg, A. E. (2006). *Constructions at work: The nature of generalization in language*. Oxford University Press.
- Goodfellow, I., Bengio, Y., & Courville, A. (2016). *Deep learning*. MIT Press.
- Hasson, U., Nastase, S. A., & Goldstein, A. (2020). Direct fit to nature: An evolutionary perspective on biological and artificial neural networks. *Neuron*, 105(3), 416–434. <https://doi.org/10.1016/j.neuron.2019.12.002>
- Jurafsky, D., & Martin, J. H. (2025). *Speech and language processing* (3rd ed., draft). Pearson.
- Kintsch, W. (1998). *Comprehension: A paradigm for cognition*. Cambridge University Press.
- Lake, B. M., Ullman, T. D., Tenenbaum, J. B., & Gershman, S. J. (2017). Building machines that learn and think like people. *Behavioral and Brain Sciences*, 40, e253. <https://doi.org/10.1017/s0140525x16001837>
- Lakoff, G. (1987). *Women, fire, and dangerous things: What categories reveal about the mind*. University of Chicago Press.
- Levelt, W. J. M. (1999). Models of word production. *Trends in Cognitive Sciences*, 3(6), 223–232.
- Linzen, T., Dupoux, E., & Goldberg, Y. (2016). Assessing the ability of LSTMs to learn syntax-sensitive dependencies. *Transactions of the Association for Computational Linguistics*, 4, 521–535. https://doi.org/10.1162/tacl_a_00115
- Mahowald, K., Ivanova, A., Blank, I., Kanwisher, N., Tenenbaum, J. B., & Fedorenko, E. (2024). Dissociating language and thought in large language models. *Trends in Cognitive Sciences*. <https://doi.org/10.1016/j.tics.2024.01.011>
- Manning, C. D., Clark, K., Hewitt, J., Khandelwal, U., & Levy, O. (2020). Emergent linguistic structure in artificial neural networks trained by self-supervision. *Proceedings of the National Academy of Sciences*, 117(48), 30046–30054. <https://doi.org/10.1073/pnas.1907367117>
- Mikolov, T., Chen, K., Corrado, G., & Dean, J. (2013). Efficient estimation of word representations in vector space. In *Proceedings of the International Conference on Learning Representations (ICLR)*.
- OpenAI. (2023). *GPT-4 technical report*. arXiv. <https://doi.org/10.48550/arXiv.2303.08774>
- Page, M. J., McKenzie, J. E., Bossuyt, P. M., Boutron, I., Hoffmann, T. C., Mulrow, C. D., et al. (2021). The PRISMA 2020 statement: An updated guideline for reporting systematic reviews. *BMJ*, 372, n71.
- Piantadosi, S. T. (2023). *Modern language models refute Chomsky's approach to language*. LingBuzz.
- Pickering, M. J., & Gambi, C. (2018). Predicting while comprehending language. *Psychological Bulletin*, 144(10), 1002–1044. <https://doi.org/10.1037/bul0000158>
- Posner, M. I., & Petersen, S. E. (1990). The attention system of the human brain. *Annual Review of Neuroscience*, 13, 25–42. <https://doi.org/10.1146/annurev.ne.13.030190.000325>
- Rumelhart, D. E., McClelland, J. L., & PDP Research Group. (1986). *Parallel distributed processing: Explorations in the microstructure of cognition* (Vols. 1–2). MIT Press.

-
- Schrimpf, M., Blank, I., Tuckute, G., Kauf, C., Hosseini, E., Kanwisher, N., ... Fedorenko, E. (2021). The neural architecture of language: Integrative modeling converges on predictive processing. *Proceedings of the National Academy of Sciences*, 118(45), e2105646118.
- Searle, J. R. (1980). Minds, brains, and programs. *Behavioral and Brain Sciences*, 3(3), 417–457. <https://doi.org/10.1017/S0140525X00005756>
- Tomasello, M. (2003). *Constructing a language: A usage-based theory of language acquisition*. Harvard University Press.
- Traxler, M. J. (2014). *Introduction to psycholinguistics: Understanding language science* (2nd ed.). Wiley-Blackwell.
- Vaswani, A., Shazeer, N., Parmar, N., Uszkoreit, J., Jones, L., Gomez, A. N., ... Polosukhin, I. (2017). Attention is all you need. *Advances in Neural Information Processing Systems*, 30, 5998–6008.