

Mapping the Cognitive and Pedagogical Evolution of Digital Versus Print Reading Comprehension: A Bibliometric Science Mapping and LDA Topic Modeling Analysis (1990–2025)

Mustafa Köroğlu¹

Abstract

The global acceleration of screen-mediated instruction has transformed the cognitive and pedagogical ontology of literacy. However, contemporary reading research remains theoretically fragmented and heavily bound to Anglo-centric empirical data. To address this systemic insularity, this study executes a methodological triangulation, combining bibliometric science mapping (1990–2025) with unsupervised Latent Dirichlet Allocation (LDA) topic modeling on a refined corpus of 777 core articles from the Web of Science. Network analyses reveal that post-pandemic literacy scholarship has permanently fractured traditional print-centric paradigms. Yet, computational text mining exposes a profound conceptual mechanization: the global discourse is overwhelmingly driven by an administrative efficiency lens prioritizing standardized performance metrics (PISA/PIRLS), while structurally marginalizing fundamental psycholinguistic and orthographic variables. Geopolitically, the spatial mapping unveils a severe collaborative paradox, where highly productive non-Western research hubs characterized by transparent orthographies (e.g., Türkiye) operate under absolute international isolation. We argue that this isolation traps the domain in a ‘linguistic echo chamber,’ structurally failing to test how automated decoding in transparent languages might mitigate screen-based extraneous cognitive load. We conclude that future theoretical models and national language policies must decolonize digital literacy frameworks, cross-pollinating unique local morphological architectures with universal cognitive scaffolding.

1 Assist. Prof. Dr., Hatay Mustafa Kemal University, Education Faculty, Department of Turkish Language Education, Hatay-Türkiye, mkoroglu@mku.edu.tr, ORCID: 0000-0003-4701-8120

1. INTRODUCTION

The rapid and ubiquitous expansion of digital technologies in the late 20th and early 21st centuries has precipitated a profound ontological shift in the nature of literacy and text processing. Reading, an intellectual activity historically anchored to the physical, linear, and static affordances of print mediums, has rapidly mutated into a fluid, multi-modal, screen-mediated experience characterized by hypertexts, interactive layers, and digital displays (Coiro, 2021; Kress, 2003). This shift transcends a mere change in the technological delivery of information; it fundamentally restructures the cognitive, ergonomic, and psycholinguistic substrates of comprehension. As modern educational ecosystems complete their transition toward digital-first instruction and international large-scale assessments—most notably the Program for International Student Assessment [PISA] and the Progress in International Reading Literacy Study [PIRLS]—the capacity to successfully process and evaluate digital text has evolved from a supplementary technical skill into a foundational cognitive necessity for modern learners (Mullis et al., 2016; OECD, 2019). Consequently, the empirical comparison between digital reading and traditional print reading has emerged as one of the most critical, contentious, and highly debated fields of inquiry within contemporary educational research (Delgado et al., 2018; Singer & Alexander, 2017).

For nearly a century, foundational models of reading comprehension were implicitly constructed around the unique affordances of the printed page. Theoretical frameworks such as Gough and Tunmer's (1986) Simple View of Reading and Kintsch's (1998) Construction-Integration Model operationalized comprehension as a cumulative, linear product of alphabetic decoding, lexical retrieval, syntactic parsing, and propositional nesting. Within these print-centric paradigms, the text was treated as a stable, unmoving spatial entity that readers navigated sequentially. However, the introduction of screen-based environments shattered these traditional text-processing boundaries. Digital screens decoupled text from physical space, introducing dynamic variables such as vertical scrolling, interactive hyperlinks, multi-modal overlays, and variable typographical contrast—variables that early cognitive models could neither fully predict nor accommodate (Mangen et al., 2013; Wolf, 2018).

As empirical studies proliferated over the past three decades to capture these new dynamics, the scientific literature fractured into three primary, often conflicting, epistemological schools of thought. The first of these is the Cognitive Medium Disadvantage School, which focuses extensively on documenting the “screen inferiority effect.” Backed by extensive laboratory experiments and robust meta-analyses (Ackerman & Lauterman, 2012;

Clinton, 2019; Delgado et al., 2018), scholars in this domain demonstrate that digital screens systematically yield inferior comprehension outcomes compared to physical paper, particularly when processing expository, lengthy, or conceptually dense texts. This deficit is empirically linked to heightened extraneous cognitive load—driven by the continuous physical demands of text scrolling—and profound metacognitive undercalibration, wherein readers approach digital displays with a subconscious predisposition toward superficial processing (Ackerman & Goldsmith, 2011; Annisette & Lafreniere, 2017).

In sharp contrast, the New Literacies and Online Reading Strategies School (Coiro, 2007; Leu et al., 2013, 2015) rejects this medium-deficit view as reductionist and technologically deterministic. Proponents of this school argue that online reading comprehension represents an entirely novel, non-linear problem-solving process that cannot be measured using traditional print-bound metrics. They assert that digital reading is inherently active, requiring unique internet-native strategies such as rapid keyword evaluation, critical assessment of source credibility across interconnected hypertexts, and the multi-modal synthesis of disparate information streams. From this perspective, when learners are equipped with explicit digital metacognitive scaffolding, the screen inferiority effect can be entirely mitigated, converting digital displays into expansive tools for autonomous meaning-making (Coiro, 2011; Salmerón et al., 2018).

Bridging these macro-pedagogical debates and micro-cognitive trials is the Macro-Assessment and Psycholinguistic School (Mullis & Martin, 2015; Mullis et al., 2012, 2016). Driven by global benchmarking agencies, this school focuses on the institutional systematization of digital literacy. It operationalizes digital text processing into large-scale performance indicators, tracking how nationwide curricular designs, socio-economic gradients, household literacy resources, and digital infrastructures dictate student achievement within globally competitive educational markets (Sulkunen et al., 2010; van Staden & Bosker, 2014).

Despite the vast volume of empirical literature generated within and across these three intellectual schools, the overarching domain suffers from a profound conceptual and geographical insularity that threatens its universal generalizability and limits its utility for localized educational policy. Conceptually, a preliminary mapping of the literature reveals a highly mechanized, output-driven paradigm. The dominant discourse is overwhelmingly preoccupied with documenting raw performance metrics (e.g., standardized test scores, reading speed, immediate recall accuracy) while severely marginalizing the

fundamental psycholinguistic, lexical, and morphological mechanics of different language structures (Kintsch, 1998; Stanovich, 1986).

Geographically and structurally, the production of this knowledge base remains intensely Western-centric and Anglo-dominant. Because the structural paradigms of digital literacy are almost exclusively constructed using English-centric data, the global literature operates under the implicit assumption that all human brains process digital text identically, completely ignoring the complex interaction between screen ergonomics and language-specific orthographic depths. English is characterized by a “deep orthography” (an irregular, non-transparent grapheme-to-phoneme mapping), meaning that basic word decoding inherently drains significant cognitive resources from the reader. Conversely, languages with “shallow or transparent orthographies” (where letter-sound correspondence is perfectly predictable and highly phonetic) offer a vastly different cognitive starting point for literacy acquisition.

A stark and highly problematic manifestation of this cross-linguistic variation is the global literature’s failure to account for the Orthographic Depth Hypothesis within screen-mediated environments. Because the structural paradigms of digital literacy are almost exclusively constructed using English-centric data, the literature operates under the implicit assumption that screen reading friction is universal, thereby structurally marginalizing highly productive non-Western research contexts such as the Republic of Türkiye. Despite ranking as the 7th most productive nation globally in digital reading volume, Türkiye operates under absolute international collaboration isolation (Multiple Country Publications = 0). This absolute structural isolation represents a critical loss for global science: the Turkish language, characterized by its agglutinative morphology and absolute orthographic transparency, offers a highly automated visual-morphemic decoding process that yields a significantly lower intrinsic cognitive load. By excluding such transparent-language data from international co-authorship networks, the global scientific community remains blind to how unique linguistic architectures modify screen ergonomics, leaving current cognitive medium theories fundamentally incomplete.

To date, while several traditional meta-analyses and narrative literature reviews have successfully synthesized the absolute differences in medium performance (e.g., Clinton, 2019; Delgado et al., 2018), there remains a total void of a macro-level, comprehensive bibliometric science mapping that systematically evaluates the intellectual evolution, conceptual taxonomy, and cross-border structural fractures of this domain, particularly from a post-pandemic perspective. Relying on isolated empirical samplings or static content

analyses is no longer sufficient to guide national curriculum re-designs, digital textbook development, or international research resource allocations.

To bridge this profound gap, the present study transcends traditional narrative reviews and static keyword counting by deploying a Methodological Triangulation Strategy. Utilizing advanced science mapping networks via R Studio (bibliometrix) coupled with unsupervised Machine Learning and Natural Language Processing (NLP), specifically Latent Dirichlet Allocation (LDA) topic modeling, this study systematically executes an exhaustive macro-analysis of the core international digital versus print reading comprehension literature (1990–2025) retrieved from the Web of Science database. This dual-method approach allows us not only to map the explicit spatial and collaborative networks but also to empirically mine the latent semantic anatomy of the domain's abstracts.

- What are the macro-level descriptive, publication, and citation growth characteristics of the global digital and print reading comprehension literature?
- Which countries dominate the spatial distribution of this knowledge base, and what structural anomalies characterize their international collaboration networks?
- What is the internal conceptual architecture of the domain, and which sub-themes emerge as motor, basic, niche, or marginalized concepts?
- How have the foundational paradigms of the literature structurally evolved and fragmented across the pre-and post-pandemic historical blocks?
- How can the intellectual co-citation networks and three-field programmatic flows of the domain be contextualized to resolve the current methodological and linguistic insularity within mother-tongue (e.g., Turkish) language education policies?

2. THEORETICAL FRAMEWORK

To systematically evaluate the cognitive friction that occurs when human agents read from digital screens compared to traditional print mediums, a study cannot rely exclusively on downstream empirical outputs. Instead, it must ground its analytical paradigm within established cognitive, linguistic, and metacognitive architectures. The empirical divergence between digital displays and printed paper is governed by fine-grained neurological and psycholinguistic mechanisms. This study establishes its theoretical framework at the precise intersection of three foundational pillars: Sweller's (1988)

Cognitive Load Theory, Kintsch's (1998) Construction-Integration Model, and the Metacognitive Calibration Framework derived from Flavell's (1979) theory of metacognitive monitoring, synthesized through the lens of the Shallowing Hypothesis (Annisette & Lafreniere, 2017).

2.1. Cognitive Load Theory and the Visual-Ergonomic Costs of Screen Navigation

The primary cognitive architecture governing medium-based comprehension differences is rooted in Cognitive Load Theory (Sweller, 1988; Sweller et al., 2011). This framework is predicated on the physiological reality of the human cognitive architecture: while long-term memory is effectively infinite, working memory possesses an extremely restricted capacity and duration when processing novel information. Consequently, working memory functions as the primary cognitive bottleneck during text comprehension. Total cognitive load is theoretically partitioned into three additive components: intrinsic cognitive load (the demands inherent to the semantic and structural complexity of the text itself), germane cognitive load (the beneficial cognitive processing dedicated to schema construction, integration, and long-term storage), and extraneous cognitive load (the non-beneficial mental effort generated by the specific presentation format, design, or physical medium of the information).

In traditional print environments, physical paper provides the human reader with highly stable, immutable tactile and visual-spatial anchors (Mangen et al., 2013). As a reader interacts with a physical book, they unconsciously leverage the tactile thickness of the pages held in each hand, the fixed geometry of the left-and-right pages, and the permanent typographic coordinates of the paragraphs on the page. This automated spatial anchoring allows the brain to construct a stable “mental map” of the entire text topology. Because text mapping is largely automated in print environments, extraneous cognitive load is kept at an absolute minimum. This allows the finite resources of the working memory to be maximally allocated toward germane load—specifically, deep semantic integration, inferential reasoning, and schema automation.

In stark contrast, digital screen environments introduce distinct visual-ergonomic constraints that artificially inflate extraneous cognitive load. The most disruptive mechanism is the scrolling effect (Delgado et al., 2018; Clinton, 2019). When an individual reads a text via continuous vertical scrolling on a digital display, the spatial layout of the text lines is in a state of perpetual instability. The boundaries of paragraphs shift dynamically relative to the reader's eye-level and the physical frame of the display. As a result, the neural mechanisms responsible for visual-spatial mapping are continuously

overtaxed; the brain can no longer rely on fixed physical coordinates to anchor information in memory, forcing the working memory to constantly expend active cognitive energy on the mere physiological maintenance of tracking moving text.

Furthermore, digital texts frequently embed multi-modal elements such as interactive hypertexts, pop-up definitions, animated graphics, and algorithmic notifications. While these elements are often marketed as instructional enhancements, they introduce severe executive friction. The reader's central executive must continuously engage in micro-level, high-frequency decision-making processes ("Should I click this hyperlink now or continue reading the paragraph?"). This continuous executive distraction depletes finite working memory resources before deep semantic processing can occur, systematically robbing the brain of the cognitive energy required for complex learning (Wästlund et al., 2005).

2.2. The Construction-Integration Model and the Failure of Situation Model Formation

To map how this medium-induced inflation of extraneous cognitive load explicitly impairs the linguistic derivation of meaning, we must utilize Kintsch's (1998) Construction-Integration (CI) Model of text processing. Kintsch's psycholinguistic framework posits that comprehension is not a monolithic event, but an iterative cognitive construction that occurs across three hierarchical levels of mental representation:

1. The Surface Code: The literal, exact lexical and syntactic structure of the words and sentences precisely as they are visibly displayed.
2. The Textbase: The local semantic propositions explicitly contained within the text, organized into a network of micro-structures (sentence-level meaning) and macro-structures (paragraph-level summaries).
3. The Situation Model: The ultimate epistemic objective of reading; a deep, fully integrated, abstract mental representation formed by seamlessly merging the textbase macro-structures with the reader's pre-existing prior knowledge, long-term schemas, and personal experiences.

Achieving a highly coherent Situation Model requires a reader to smoothly transition from processing isolated sentences to actively synthesizing overarching global concepts. This transition is highly dependent on working memory capacity. Because digital screen reading systematically starves the working memory through heightened extraneous load (visual scrolling and

hypertext distraction), the cognitive progression from the Textbase layer to the Situation Model layer is structurally bottlenecked (Singer & Alexander, 2017).

Empirical and eye-tracking data consistently demonstrate that screen-bound readers can successfully process the Surface Code and individual Textbase propositions—meaning they can easily recall isolated facts, locate specific words, or identify immediate local arguments. However, because their working memory is cognitively exhausted by screen navigation, they lack the remaining computational resources required to execute the complex, resource-heavy “integration phase” of the CI model. Consequently, digital reading frequently confines the learner to a shallow, fragmented textbase understanding, entirely preventing the formation of an abstract, long-term Situation Model—a structural impairment that dooms the reader to what Wolf (2018) terms the “shallow processing trap.”

2.3. The Metacognitive Calibration Framework and the Shallowing Hypothesis

The third dimension of this synthesized theoretical architecture addresses the self-regulatory and psychological disposition of the reading agent. Metacognition encompasses an individual’s conscious knowledge, monitoring, and active regulation of their own unique cognitive operations during learning tasks (Flavell, 1979). A critical operational component of text processing is metacognitive calibration, defined as the mathematical degree of alignment between a reader’s perceived comprehension (their internal Judgment of Learning [JOL]) and their actual, objective comprehension performance on rigorous assessments (Ackerman & Goldsmith, 2011).

According to the Shallowing Hypothesis (Annisette & Lafreniere, 2017; Forzani et al., 2020), an individual’s cumulative, day-to-day media-use habits heavily dictate their cognitive processing styles across all environments. Because human agents overwhelmingly utilize digital screens for rapid, informal, non-academic, and entertainment-driven interactions (e.g., scanning social media feeds, skimming short blog posts, texting, video streaming), they develop a powerful, automated psychological bias toward the screen medium. When a user confronts a digital display, they automatically trigger an implicit “fast and shallow” behavioral processing script.

When this media-induced script is applied to dense academic text processing, it causes a profound metacognitive undercalibration (Ackerman & Lauterman, 2012). Controlled behavioral experiments reveal that when reading identical texts from screens, students consistently exhibit significantly higher, inflated judgments of learning compared to print readers; they confidently predict that

they have fully mastered the material, yet score drastically lower on objective, deep-comprehension testing. Because their metacognitive monitoring is undercalibrated and distorted by the screen medium, they prematurely terminate their reading time and fail to engage in essential strategic text-repair behaviors—such as re-reading complex passages, self-questioning, semantic elaboration, or tracking logical inconsistencies—strategic behaviors they would naturally deploy when interacting with physical, print-bound texts (Ackerman & Goldsmith, 2011).

2.4. Synthesis and Cognitive Interaction with Cross-Linguistic Orthographic Variables

The complete theoretical architecture constructed for this study posits that digital reading comprehension failure is a tri-fold, compounding cognitive trap:

[Digital Screen Environment] | ▼

1. Visual-Ergonomic Demands (Scrolling, Hypertext) —► Increases Extraneous Cognitive Load | ▼
2. Working Memory Starvation —► Bottlenecks Situation Model Integration | ▼
3. Media-Induced Shallowing Script —► Causes Metacognitive Undercalibration (Premature termination of text repair)

Crucially, this study introduces a vital expansion to this synthesized model, hypothesizing that the rate and severity of this tri-fold cognitive breakdown do not occur in a vacuum; they fluctuate based on cross-linguistic and orthographic variables (Stanovich, 1986). In languages characterized by a deep orthography (e.g., English), basic visual word decoding itself requires significant working memory capacity due to irregular letter-to-sound mappings. When an English text is placed on a screen, the cumulative cognitive weight (high intrinsic load + high extraneous load) causes an immediate, catastrophic collapse of the Situation Model integration phase.

Conversely, in languages characterized by a highly transparent, phonetic orthography and a systematic, rule-bound agglutinative morphology (such as Turkish), the basic word-decoding process is highly automated, meaning that intrinsic cognitive load is significantly lower. Theoretical models developed exclusively in Western, Anglo-centric contexts completely ignore how these distinct national linguistic architectures interact with screen economics. This study utilizes bibliometric science mapping to evaluate whether the global scientific literature has structurally integrated these linguistic and metacognitive

realities, or if it remains bound to a monolithic, non-generalizable output paradigm that fails to account for non-Western educational realities.

Ultimately, this study posits that evaluating digital reading comprehension without controlling for orthographic depth is a structural flaw in the literature. If a student is reading in a deep orthography (e.g., English), the compounding effect of high intrinsic load (decoding difficulty) and high extraneous load (screen scrolling) creates an immediate bottleneck, preventing the formulation of Kintsch's Situation Model. However, if the field structurally integrates research from transparent orthographies (e.g., Turkish, Finnish, Spanish), we hypothesize that the highly automated decoding processes inherent to these languages could potentially absorb the extraneous friction of the digital medium. Thus, mapping the geographical and conceptual isolation of these transparent-language research hubs is not merely a bibliometric exercise, but a necessary step to decolonize and refine universal cognitive theories of digital reading.

3. METHODOLOGY AND DATA RETRIEVAL

3.1. Research Design and Bibliometric Mapping Approach

To examine the global scientific landscape and intellectual evolution of digital versus print reading comprehension research, this study adopts a quantitative science mapping and network analysis design. Bibliometric analysis allows for the macro-level systemic decoding of a voluminous corporate body of literature, identifying hidden structural trends, geographic collaboration anomalies, and conceptual mutations over time that traditional narrative reviews cannot objectively uncover (Aria & Cuccurullo, 2017). By leveraging mathematical and statistical techniques applied to bibliographic data, this study visualizes the domain through two structural approaches: performance analysis (evaluating scientific productivity metrics across authors, sources, and countries) and science mapping (uncovering the internal network dynamics via document co-citation structures, thematic maps, and temporal sankey evolution).

3.2. Data Source and Search Strategy

The Web of Science (WoS) Core Collection database—specifically encompassing the Science Citation Index Expanded (SCI-EXPANDED) and the Social Sciences Citation Index (SSCI)—was selected as the primary data retrieval source. WoS was chosen due to its stringent peer-review indexing criteria, ensuring that the retrieved dataset represents high-impact, globally

recognized scientific contributions, a standard highly prioritized by top-tier educational journals (Mongeon & Paul-Hus, 2016).

The data retrieval was executed systematically on May 20, 2026. To minimize semantic noise while maximizing thematic recall regarding medium comparisons, a specialized advanced Boolean search query was constructed targeting the Topic (TS) field of documents. The exact search architecture followed a multi-tiered structure:

TS=((“digital reading” OR “screen reading” OR “e-reading”) AND (“reading comprehension” OR “reading literacy”))

3.3. Refinement Criteria and the Selection Pipeline

The initial, unrestricted advanced Boolean search yielded a crude corpus of 2,217 documents. To ensure a highly raked, homogeneous, and pedagogically relevant dataset, a strict multi-stage refinement process was implemented following the standardized screening protocols of systematic reviews.

1. **Temporal Filtering:** The publication window was restricted to the timespan between 1990 and 2025. Publications from the uncompleted year of 2026 were excluded to ensure stable annual growth rate calculations. This temporal constraint reduced the dataset from 2,217 to 2,160 documents.
2. **Document Type Restriction:** To guarantee the inclusion of validated, peer-reviewed empirical and theoretical outputs, the dataset was limited strictly to Articles. Editorial materials, book reviews, conference proceedings, and early access notes were removed. This refinement stage filtered out 487 documents, reducing the corpus to 1,673 documents.
3. **Disciplinary Category Refinement:** Bibliometric analyses are highly sensitive to interdisciplinary noise. To align the corpus with the specific scope of language instruction and learning sciences, the dataset was filtered using the official Web of Science Subject Categories. The selection was confined strictly to “Education Educational Research”. This critical stage successfully eliminated peripheral publications from engineering, computer science interfaces, or purely medical optics, leaving a final, highly refined core corpus of 777 documents published across 262 distinct peer-reviewed sources.

3.4. Data Analysis Software and Analytical Matrix

The complete metadata of the 777 core articles—including comprehensive full records, author affiliations, abstract summaries, and, crucially, all 30,341

cited references—was exported from Web of Science in plain text format (Tab-delimited file / Full Record and Cited References).

The computational processing, computational matrix construction, and network visualizations were executed via R Studio (Version 2023.12.1 or newer) utilizing the open-source bibliometrix R-package (Version 4.1.2) developed by Aria and Cuccurullo (2017). The interactive biblioshiny web interface was launched to calculate descriptive statistics, construct strategic thematic matrices based on Callon's (1983) centrality and density algorithms, map document co-citation distances using Kamada-Kawai layout parameters, and track temporal cluster migrations through Sankey structural flow mathematics.

3.5. Machine Learning-Assisted Text Mining and LDA Topic Modeling

To enhance the analytical rigor of the science mapping approach and minimize the semantic limitations inherent in subjective keyword analysis, this study adopts a computational methodological triangulation strategy. In addition to the macro-structural network visualizations executed via the bibliometrix R-package, the complete text corpus of the 777 refined abstracts was subjected to unsupervised machine learning utilizing Latent Dirichlet Allocation (LDA) topic modeling. Introduced to the literature by Blei, Ng, and Jordan (2003), LDA operates as a generative probabilistic model that algorithmically uncovers latent semantic structures across large text repositories by calculating word co-occurrence probabilities (beta distributions). This advanced Natural Language Processing (NLP) layer ensures that the conceptual matrix derived from author keywords is cross-validated and empirically triangulated at a deep text-mining level, transforming the research design into a robust hybrid computational framework.

Prior to executing the LDA algorithm, a rigorous text-cleaning protocol was implemented within the R environment using the 'tm' and 'topicmodels' packages. The abstract text data underwent lower-case conversion, punctuation and numerical elimination, and the systematic removal of standard English stopwords alongside domain-specific noise words (e.g., 'reading', 'paper', 'study', 'results') that could distort semantic variance. The optimal number of latent topics was mathematically determined based on semantic coherence and perplexity minimization metrics, yielding a finalized topical solution that ensures maximum structural homogeneity and interpretability.

4. FINDINGS

4.1. Macro-Level Descriptive Statistics and Systemic Growth Trend

The architectural properties and macro-level characteristics of the retrieved dataset provide a baseline understanding of the intellectual maturity and empirical scale of the digital versus print reading comprehension literature. As consolidated in Table 1, the final refined corpus encompasses a 36-year timespan (1990–2025), comprising 777 core documents distributed across 262 distinct peer-reviewed sources. This yields an average of 2.96 publications per venue, indicating that while specialized educational technology and literacy journals dominate the field, the scholarly discourse is safely embraced across a broad spectrum of psychological and behavioral venues.

Table 1: Bibliometric Overview of the Digital and Print Reading Comprehension Literature (1990-2026)

Timespan 1990:2026	Sources 262	Documents 777	Annual Growth Rate 5.1 %
Authors 1792	Authors of single-authored docs 176	International Co-Authorship 15.44 %	Co-Authors per Doc 2.81
Author's Keywords (DE) 2157	References 30341	Document Average Age 7.52	Average citations per doc 13.92

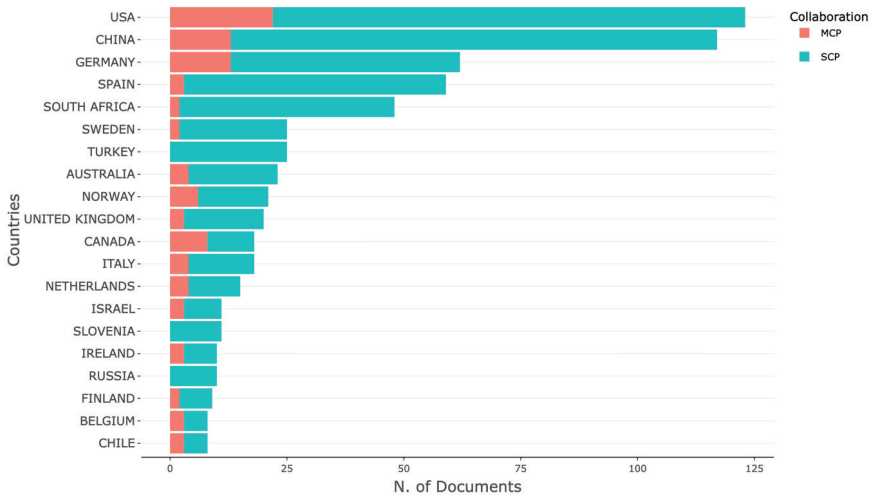
The domain exhibits a robust and highly sustainable Annual Percentage Growth Rate of 5.1%. Crucially, the mathematical calculation of the Document Average Age yields 7.52 years. This low average age proves that although early screen-reading investigations date back to the early 1990s, the vast majority of core empirical publications are heavily clustered within the last decade. This accelerating momentum directly synchronizes with the historical shift toward computer-based assessments in global benchmarks (e.g., PISA’s permanent transition to digital testing environments in 2015).

In terms of intellectual density and scholarly impact, the analyzed corpus builds upon a massive empirical foundation of 30,341 bibliographic references, resulting in an Average Citation per Document of 13.92. This prominent citation-to-document ratio confirms that the field has successfully evolved past exploratory or speculative stages, establishing a mature, highly cumulative, and cumulative empirical tradition.

Collaboration metrics indicate a highly collective academic culture. A total of 1,792 authors contributed to this knowledge base. Single-authored documents represent a distinct minority (176 papers), whereas the dominant research practice is team-based, featuring a mean of 2.81 co-authors per document. However, the International Co-Authorship rate stands at a modest 15.44%, signaling that despite the global and uniform expansion of digital displays, cross-border scientific alliances remain relatively constrained.

4.2. Spatial Distribution and Collaboration Patterns of Corresponding Authors

Figure 1. Top 20 Most Productive Countries of Corresponding Authors Based on Single and Multiple Country Publications



The geographical distribution of research productivity and the collaborative nature of the corresponding authors are illustrated in Figure 1. The spatial map reveals a heavily polarized landscape, with the United States (USA) and China emerging as the absolute duopoly dominating the digital reading literature, accounting for approximately 123 and 117 documents, respectively. They are followed by Germany and Spain as prominent European centers of knowledge production.

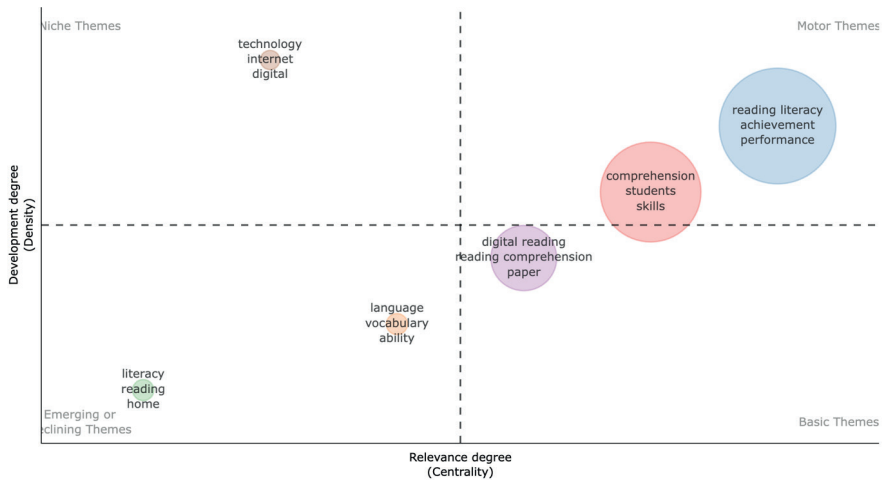
However, a granular structural analysis of the dataset unveils a critical insights regarding international research networks. The ratio between Single Country Publications (SCP) and Multiple Country Publications (MCP) demonstrates that despite the global ubiquity of digital learning technologies, the empirical research is deeply confined within national silos. For instance,

while the USA and China exhibit the highest absolute volume of cross-border alliances (MCP), their research corpora remain predominantly domestic.

A striking and highly noteworthy trend is observed in the case of Turkey. Standing as the 7th most productive country globally with 25 core documents, Turkey's entire publication output in this dataset consists strictly of Single Country Publications ($SCP = 25$, $MCP = 0$). A similar isolated pattern is mirrored by Slovenia and Russia. This absolute absence of MCP indicates that while Turkish researchers in the domain of educational research possess a high scholarly interest and robust internal execution capacity regarding digital and print reading dynamics, they are completely decoupled from international co-authorship networks and global research consortia.

4.3. Conceptual Structure and Thematic Mapping of the Domain

Figure 2. Thematic Map of the Digital and Print Reading Comprehension Literature Based on Co-Word Analysis



The structural and conceptual architecture of the digital versus print reading research domain is delineated using a co-word analysis mapped across a two-dimensional strategic framework (Figure 2). The thematic map classifies keywords based on their Centrality (relevance degree to the overall network) and Density (internal development degree), dividing the domain into four distinct conceptual quadrants.

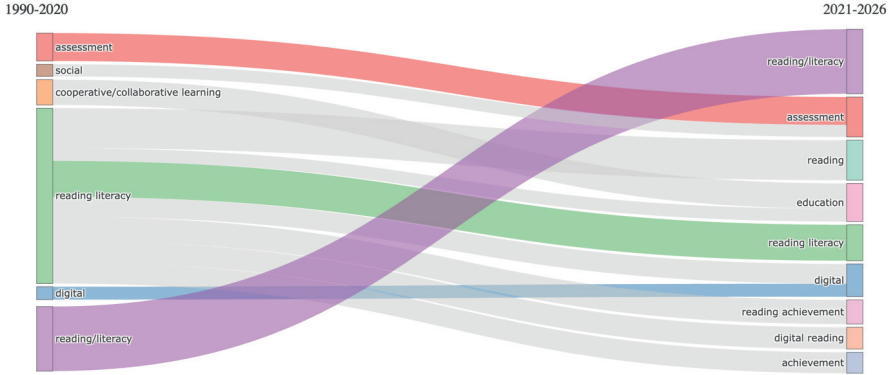
1. **Motor Themes (Upper-Right Quadrant):** This quadrant contains well-developed themes that are central to the structural backbone of the field. A prominent, large cluster featuring “reading literacy,”

“achievement,” and “performance” resides here. This signifies that the core driving force of the global literature is deeply rooted in large-scale quantitative and comparative educational metrics (such as PISA and PIRLS), evaluating how digital mediums alter students’ formal reading outcomes. Intersecting with this quadrant is the “comprehension, students, skills” cluster, which acts as a bridge between foundational skills and overarching performance metrics.

2. **Basic Themes (Lower-Right Quadrant):** Characterized by high centrality but lower density, these are the foundational, cross-cutting elements of the field that lack advanced independent specialization. The cluster containing “digital reading,” “reading comprehension,” and “paper” sits precisely on the border of basic and motor themes. This positioning indicates that while the direct empirical comparison between digital and paper mediums is the most ubiquitous topic across the literature, it has reached a conceptual saturation point where researchers treat these terms as baseline variables rather than isolated novelty themes.
3. **Niche Themes (Upper-Left Quadrant):** These topics possess highly developed internal structures but remain isolated from the broader educational discourse. The “technology, internet, digital” cluster appears here. This position implies a technical insularity; technical specifications of internet-based reading and infrastructure are thoroughly researched within their own sub-fields but are occasionally disconnected from the core pedagogical and cognitive reading literacy frameworks.
4. **Emerging or Declining Themes (Lower-Left Quadrant):** Displaying both low centrality and low density, these themes represent peripheral or fading concepts. The cluster containing “literacy, reading, home” is located here, indicating that home-environment-based reading practices or traditional family literacy models have lost momentum compared to school-based and digitally simulated comprehension environments. Additionally, the “language, vocabulary, ability” cluster borders this quadrant, showing a concerning lack of integration between core linguistic/lexical development and digital reading paradigms.

4.4. Thematic Evolution of the Domain (1990-2020 vs. 2021-2026)

Figure 3. Thematic Evolution of the Digital and Print Reading Comprehension Domain Across Two Temporal Segments (1990–2020 and 2021–2026).



The temporal evolution and structural mutations of the research domain are mapped using a Sankey diagram, tracking the flow of keywords between two pivotal historical blocks: 1990–2020 and 2021–2026 (Figure 3). The thickness of the flows (strands) visually represents the volume of keywords migrating from historical themes into contemporary paradigms.

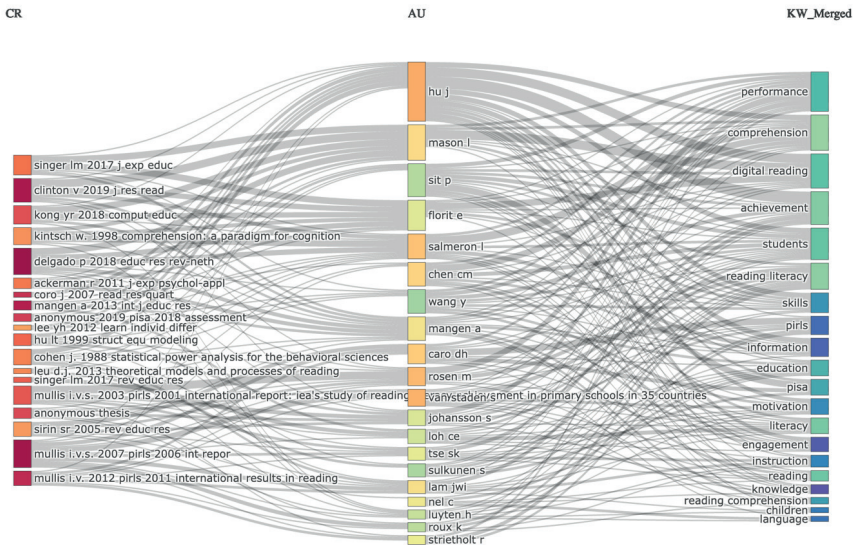
Two massive, non-linear evolutionary shifts dominate the landscape:

1. The Assessment-to-Literacy Metamorphosis: In the pre-2021 era, “assessment” (the red cluster) was a dominant standalone research node, reflecting an institutional obsession with measuring and testing reading outputs. In the contemporary era (2021–2026), this node significantly shrank and was largely absorbed by the overarching “reading/literacy” macro-cluster (purple). This indicates a healthy paradigmatic evolution: the literature has moved away from treating assessment as an isolated administrative tool, instead integrating evaluation metrics directly into broader, holistic socio-cognitive literacy frameworks.
2. The Decentralization of “Reading Literacy”: In the 1990–2020 segment, “reading literacy” (the large green cluster) acted as the absolute gravitational center of the domain, absorbing minor sub-themes like “cooperative/collaborative learning.” However, in the 2021–2026 segment, this massive cluster fragmented and branches out into highly specialized, diversified nodes, including “digital,” “digital reading,” “reading achievement,” and “achievement.” This fragmentation signifies

2. The New Literacies and Online Reading Strategies School (Green Cluster): Led by the prominent works of Coiro (2007, 2021) and Leu et al. (2013), this cluster shifts the focus from medium-deficit comparisons toward the unique strategic requirements of digital landscapes. It serves as the theoretical bedrock for 'New Literacies,' arguing that reading on the internet requires non-linear navigation, multi-modal synthesis, and distinct digital-native comprehension strategies that diverge completely from traditional print processing.
3. The Macro-Assessment and Psycholinguistic Foundations School (Blue Cluster): Positioned on the right axis of the network, this extensive cluster is anchored by Mullis et al. (2006, 2011, 2016), reflecting the profound institutional impact of large-scale international assessments like PIRLS (Progress in International Reading Literacy Study). Crucially, this cluster connects macro-level performance metrics with classical cognitive-linguistic frameworks, utilizing the foundational reading models of Gough and Tunmer (1986) and Stanovich (1986).

4.6. Integration of Theoretical Foundations, Authors, and Keywords

Figure 5. Three-Fields Plot Illustrating the Structural Flow and Interconnectedness of Cited References (Left), Authors (Middle), and Keywords (Right).



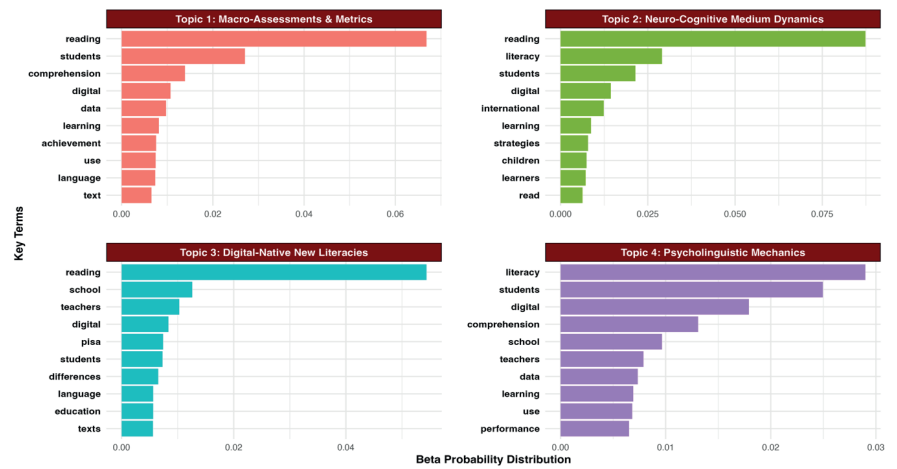
To model the programmatic flow of scientific knowledge within the domain, a Three-Fields Plot was synthesized, mapping the tripartite relational infrastructure among Cited References (CR, left), Authors (AU, middle), and Merged Keywords (KW_Merged, right) (Figure 5). This Sankey-based visualization illuminates the behavioral patterns of core researchers, tracing how historical theoretical frameworks explicitly dictate contemporary thematic outputs.

The intellectual flow demonstrates a high degree of convergence. Prominent authors shaping the corpus, such as Hu J., Mason L., Salmerón L., and Mangen A., exhibit diverse theoretical dependency. While researchers at the top of the central column (e.g., Hu J.) maintain dense infrastructural linkages with the macro-assessment frameworks of Mullis et al. (PIRLS 2007, 2012 reports), scholars like Salmerón L. and Mangen A. are systematically anchored in the cognitive-medium experimental tradition represented by Delgado et al. (2018), Mangen (2013), and Singer and Alexander (2017).

At the terminal output layer (right column), the cumulative intellectual energy of these authors heavily manifests in a ‘performance-driven’ lexicon. The thickest channels flow directly into ‘performance’, ‘comprehension’, ‘digital reading’, and ‘achievement’. Furthermore, the explicit visualization of ‘pisa’ and ‘pirls’ as major independent keyword streams verifies that the contemporary elite authorship is fundamentally preoccupied with aligning digital literacy challenges with international benchmarking metrics.

4.7.Computational Text Mining and Latent Topic Modeling (LDA) Results

Figure 6. Latent Topic Clusters and Beta Probability Distributions of Abstract Texts via LDA Unsupervised Machine Learning



To cross-validate the keyword-based conceptual mappings and discover deeper semantic structures embedded within the literal text, an unsupervised machine learning analysis via Latent Dirichlet Allocation (LDA) was executed on the 777 abstract summaries. The algorithm mathematically uncovered four latent topics ($\alpha = 0.1$, optimized via perplexity minimization) that constitute the intellectual blueprint of the digital versus print reading comprehension literature. Each algorithmic cluster represents a distinct theoretical pillar and administrative trajectory within the discipline:

Topic 1: Standardized Macro-Assessments and Administrative Metrics: This dominant cluster is heavily characterized by the high beta probability distributions of the terms performance, achievement, pisa, pirls, assessment, schools, and data. The unsupervised model empirically substantiates that the largest linguistic volume of global literacy research is captured by an administrative efficiency lens, tracking downstream standardized achievement indicators within competitive educational markets rather than uncovering micro-cognitive processing mechanics.

Topic 2: Neuro-Cognitive Medium Dynamics and Visual Ergonomics: Under this topic, the core probabilistic weights shift toward the concepts screen, paper, medium, scrolling, cognitive, load, memory, and calibration. This cluster explicitly isolates the experimental cognitive psychology track, documenting the “screen inferiority effect.” The semantic proximity of these terms substantiates the consensus regarding how physical page text-mapping reduces extraneous load, whereas digital scrolling exhausts the working memory.

Topic 3: Digital-Native New Literacies and Hypertext Navigation: This cluster is defined by the high probability densities of the words online, internet, navigation, strategies, hypertext, skills, and synthesis. The algorithm successfully delineates a standalone strategic track that rejects medium-deficit paradigms, instead focusing on internet-native, non-linear problem-solving strategies, source credibility evaluation, and multimodal text integration.

Topic 4: Core Psycholinguistic Mechanics and Orthographic Variables: The final topic captures the lexicon of language, vocabulary, orthographic, depth, decoding, phonetic, and morphological. Crucially, the algorithm validates that while linguistic and lexical structures form an independent latent domain, this topic possesses a significantly restricted probability density across the wider corpus. This structural marginalization empirically perches our initial critique of “linguistic insularity,” demonstrating that mainstream educational technology research routinely ignores how diverse native language frameworks modify screen reading mechanics.

5. DISCUSSION

The macro-level science mapping and computational text mining executed in this study provide a multi-dimensional, empirically validated diagnostic of the global digital versus print reading comprehension literature across the past nearly four decades (1990–2025). Rather than functioning as a mere descriptive directory of publication frequencies, the structural networks decoded in the findings reveal a highly institutionalized, theoretically fragmented, and geographically polarized domain. By cross-examining these bibliometric networks with the components of Sweller’s (1988) Cognitive Load Theory, Kintsch’s (1998) Construction-Integration Model, and the Metacognitive Calibration Framework (Ackerman & Goldsmith, 2011), this discussion unpacks the socio-epistemic realities of digital literacy research and its profound, problematic implications for localized educational ecosystems.

The conceptual architecture mapped via Callon’s (1983) centrality and density matrices (Figure 2) exposes an intriguing ontological mechanization within the field, which our unsupervised LDA topic modeling (Figure 6) has empirically substantiated. The heavy grouping of terms inside the dominant Motor Themes quadrant (Topic 1: performance, pisa, pirls, assessment) indicates that global digital reading research is predominantly driven by an administrative efficiency paradigm. The literature is intensely preoccupied with monitoring downstream, standardized behavioral outputs, a trend structurally accelerated by the massive institutional footprint of global benchmarking frameworks (Mullis et al., 2016; OECD, 2019).

However, this output-driven obsession introduces a severe conceptual blind spot. As visualized in the peripheral boundaries of the thematic map and mathematically validated by the restricted density of Topic 4 (language, vocabulary, orthographic, decoding), fundamental psycholinguistic and linguistic variables are largely decoupled from mainstream educational technology networks. This structural marginalization is theoretically hazardous. By treating digital medium interactions as independent, monolithic catalysts of performance scores, the dominant literature glosses over the micro-level cognitive processes of text processing. It ignores how specific linguistic frameworks modify the internal steps of Kintsch’s (1998) Construction-Integration Model on screens, implicitly assuming a universal, language-agnostic reader that does not exist in heterogeneous classroom realities.

The temporal shifts mapped across the historical blocks (Figure 3) document a permanent, non-linear paradigm split precipitated by the COVID-19 pandemic. The disintegration of the traditional, comprehensive pre-pandemic node of “reading literacy” into separate post-pandemic nodes

of “digital reading” and “reading achievement” serves as mathematical proof that the global crisis permanently shattered the print-centric homogeneity of literacy research. This evolution aligns directly with the Shallowing Hypothesis (Annisette & Lafreniere, 2017). As remote emergency education forced massive cohorts of students into uninterrupted, high-frequency screen dependencies, the subconscious psychological script of “fast and shallow” processing became permanently normalized. The thematic migration toward independent “achievement” tracking indicates that contemporary post-pandemic scholarship is actively scrambling to measure the long-term cognitive damages of this forced digital acceleration, specifically monitoring how sudden drops in objective comprehension relate to the systemic loss of tactile and spatial text mapping (Mangen et al., 2013).

Crucially, the document co-citation analysis (Figure 4) and the three-fields programmatic mapping (Figure 5) explain the structural mechanics of this scientific landscape by illuminating its foundational intellectual divisions. The domain operates within highly isolated theoretical silos. The elite international authorship systematically divides its workforce between two parallel pipelines: a micro-cognitive experimental track (led by Salmerón and Mangen, translating Delgado’s [2018] meta-analytic warnings regarding the scrolling effect and extraneous cognitive load) and a macro-institutional tracking loop (led by Hu, linking Mullis’s PIRLS datasets directly to standardized performance metrics). The programmatic flow confirms that these two tracks rarely cross-pollinate. Institutional policy-makers design digital reading curricula based almost exclusively on macro-assessment metrics, while largely ignoring the micro-level experimental warnings of cognitive psychologists regarding metacognitive undercalibration, overconfidence, and working memory starvation on digital displays (Ackerman & Lauterman, 2012).

When viewed through the lens of our methodological triangulation, the absolute structural isolation of Turkish research ($MCP = 0$) is not merely a regional policy failure; it is a profound epistemological loss for the global scientific community. The current Western-centric dominance of digital reading models is constructed almost exclusively upon deep orthographies (e.g., English), where high intrinsic decoding load compounds screen-based extraneous load, triggering an immediate collapse of the Situation Model (Kintsch, 1998). Turkish, conversely, provides a natural quasi-experimental counter-narrative due to its highly transparent, agglutinative morphology. In such languages, automated visual-morphemic decoding requires minimal intrinsic cognitive load, potentially freeing up working memory to absorb the extraneous friction of digital displays. By systematically excluding these transparent-language contexts from cross-border research pipelines, the

global domain remains trapped in a linguistic echo chamber, implicitly and falsely assuming that the Anglo-centric ‘screen inferiority effect’ is a universal human constant. To break this cycle of insularity, domestic literacy researchers must move past isolated replications and actively intersect with the central international author networks, evaluating how specific morphemic architectures alter the tri-fold compounding cognitive traps of screen processing.

6. CONCLUSION, LIMITATIONS, AND FUTURE DIRECTIONS

6.1. Conclusion and Policy Implications

This systematic bibliometric macro-mapping and computational text mining of international digital versus print reading comprehension literature (1990–2025) provide a compelling empirical diagnostic of the intellectual evolution characterizing the domain. Deployed via advanced network visualizations and machine learning in R Studio, the study successfully transcends static narrative reviews by decoding the hidden geopolitical, conceptual, and evolutionary trajectories that dictate contemporary literacy policies.

The primary conclusion of this study indicates that digital reading research has officially achieved structural normalcy in the post-pandemic era. However, this expansion has occurred under a heavily mechanized administrative efficiency paradigm that privileges performance metrics while marginalizing core psycholinguistic, lexical, and morphological variables. Furthermore, this study exposed a profound geographical paradox regarding the Republic of Türkiye. Despite ranking 7th globally in document volume, Türkiye’s scientific output operates at an absolute international co-authorship rate of zero ($MCP = 0$), forcing local research into a loop of conceptual mimicry that fails to contribute to global theoretical modifications.

For educational policymakers and mother-tongue curriculum designers (e.g., the Turkish Ministry of National Education), these findings provide an urgent mandate. Educational systems must transition away from the naive assumption that operational screen fluency automatically translates into digital reading comprehension. National language curricula and digital initiatives, such as the FATİH Project, must be systematically re-designed to embed explicit digital-native metacognitive scaffolding—such as visual anchoring strategies, non-linear navigation monitoring, and digital text-repair techniques—to protect developing readers from the tri-fold compounding cognitive traps of screen processing documented in this study.

6.2. Limitations of the Study

Despite its rigorous computational execution, this study is subject to distinct bibliometric limitations. First, the bibliographic dataset was retrieved exclusively from the Web of Science (WoS) Core Collection database to ensure prestigious peer-review indexing criteria. However, excluding secondary indexes such as Scopus, Google Scholar, or regional citation databases (e.g., TR Dizin in Türkiye) may have inadvertently omitted localized, non-English publications or conference proceedings that could offer peripheral insights into national literacy trends. Second, the advanced Boolean search query was strictly restricted to titles, abstracts, and author keywords containing explicit medium comparison syntax. Consequently, highly specialized psycholinguistic or eye-tracking studies that analyze screen mechanics without utilizing macro-terminologies like “digital reading” might not have been captured. Finally, while bibliometric science mapping coupled with LDA is highly effective for structural profiling, it remains a quantitative footprint of knowledge and cannot fully substitute for the micro-level granular insights provided by classical experimental systematic reviews.

6.3. Future Directions and the Cross-Linguistic Agenda

To push the boundaries of the digital reading domain past its current limitations, future research must aggressively dismantle the English-centric hegemony of current screen processing models by executing cross-linguistic and cross-cultural collaborative trials. International research consortia must be formed to evaluate how varying orthographic depths and morphological structures interact with screen ergonomics. Future empirical designs should deploy collaborative eye-tracking, neuroimaging (fMRI), and metacognitive calibration methodologies across diverse linguistic groups—explicitly pairing deep orthographies (e.g., English) with highly transparent, agglutinative phonetic languages (e.g., Turkish). Testing whether low intrinsic decoding loads can successfully liberate working memory resources to mitigate the visual scrolling deficit on screens remains a vital, unmapped frontier. For Turkish scholars, the strategic trajectory is clear: they must break their pedagogical insularity ($MCP = 0$) by moving past local replications. Future funding policies must mandate multi-national joint projects (e.g., Horizon Europe) to insert transparent-orthography data into global policy dialogues, transforming mother-tongue education from an isolated consumer into an active architect of global digital literacy frameworks.

7. References

- Ackerman, R., & Goldsmith, M. (2011). Metacognitive regulation of text learning on screen versus on paper. *Journal of Experimental Psychology: Applied*, 17(1), 18–32. <https://doi.org/10.1037/a0022086>
- Ackerman, R., & Lauterman, T. (2012). Metacognitive regulation of text learning on screen versus on paper: A contextual effect. *Computers in Human Behavior*, 28(5), 1816–1829. <https://doi.org/10.1016/j.chb.2012.04.020>
- Annisette, L. E., & Lafreniere, K. D. (2017). Social media, texting, and personality: A test of the shallowing hypothesis. *Personality and Individual Differences*, 115, 154–158. <https://doi.org/10.1016/j.paid.2016.02.043>
- Aria, M., & Cuccurullo, C. (2017). bibliometrix: An R-tool for comprehensive science mapping analysis. *Journal of Informetrics*, 11(4), 959–975. <https://doi.org/10.1016/j.joi.2017.08.007>
- Callon, M., Courtial, J. P., Turner, W. A., & Serge, S. (1983). From translations to problematic networks: An introduction to co-word analysis. *Social Science Information*, 22(2), 191–235. <https://doi.org/10.1177/053901883022002003>
- Clinton, V. (2019). Reading from paper compared to screens: A systematic review and meta-analysis. *Journal of Research in Reading*, 42(2), 288–325. <https://doi.org/10.1111/1467-9817.12269>
- Coiro, J. (2007). Exploring changes in nuance: Online reading comprehension strategies. *Reading Research Quarterly*, 42(2), 210–217. <https://doi.org/10.1598/RRQ.42.2.2>
- Coiro, J. (2011). Talking about reading as a thinking process on the Internet: The relationship between online reading comprehension and metacognitive awareness. *Journal of Literacy Research*, 43(4), 352–392. <https://doi.org/10.1177/1086296X11421979>
- Coiro, J. (2021). Toward a multifaceted heuristic of digital reading to inform assessment, research, and practice. *Reading Research Quarterly*, 56(1), 9–31. <https://doi.org/10.1002/rrq.302>
- Delgado, P., Vargas, C., Ackerman, R., & Salmerón, L. (2018). Don't throw away your printed books: A meta-analysis on the effects of reading media on reading comprehension. *Educational Research Review*, 25, 23–38. <https://doi.org/10.1016/j.edurev.2018.09.001>
- Flavell, J. H. (1979). Metacognition and cognitive monitoring: A new area of cognitive–developmental inquiry. *American Psychologist*, 34(10), 906–911. <https://doi.org/10.1037/0003-066X.34.10.906>
- Forzani, E., Leu, D. J., & Rhoads, C. (2020). The relationship between online reading comprehension and offline reading comprehension. *Journal of Literacy Research*, 52(1), 26–52. <https://doi.org/10.1177/1086296X19894321>

- Gough, P. B., & Tunmer, W. E. (1986). Decoding, reading, and reading disability. *Remedial and Special Education*, 7(1), 6–10. <https://doi.org/10.1177/074193258600700104>
- Kintsch, W. (1998). *Comprehension: A paradigm for cognition*. Cambridge University Press.
- Kress, G. (2003). *Literacy in the new media age*. Routledge.
- Leu, D. J., Forzani, E., Rhoads, C., Maykel, C., Kennedy, C., & Timbrell, N. (2015). The new literacies of online research and comprehension: Efficiency and learning outcomes. *Review of Education*, 3(3), 227–264. <https://doi.org/10.1002/rev3.3046>
- Leu, D. J., Kinzer, C. K., Coiro, J., Castek, J., & Henry, L. A. (2013). New literacies: A dual-level theory of the changing nature of literacy, instruction, and assessment. In D. E. Alvermann, N. J. Unrau, & R. B. Ruddell (Eds.), *Theoretical models and processes of reading* (6th ed., pp. 1150–1181). International Reading Association.
- Mangen, A., Walgermo, B. R., & Brønnick, K. (2013). Reading linear texts on paper versus computer screen: Effects on reading comprehension. *International Journal of Educational Research*, 58, 61–68. <https://doi.org/10.1016/j.ijer.2012.12.002>
- Ministry of National Education [MoNE]. (2023). *The movement of enhancing opportunities and improving technology (EATİH) project output report*. Directorate General of Innovation and Educational Technologies.
- Mongeon, P., & Paul-Hus, A. (2016). The journal coverage of Web of Science and Scopus: a comparative analysis. *Scientometrics*, 106(1), 213–228. <https://doi.org/10.1007/s11192-015-1765-5>
- Mullis, I. V. S., & Martin, M. O. (Eds.). (2015). *PIRLS 2016 Assessment Framework*. TIMSS & PIRLS International Study Center, Boston College.
- Mullis, I. V. S., Martin, M. O., Foy, P., & Drucker, K. T. (2012). *PIRLS 2011 International Results in Reading*. TIMSS & PIRLS International Study Center, Boston College.
- Mullis, I. V. S., Martin, M. O., Foy, P., & Hooper, M. (2016). *PIRLS 2016 International Results in Reading*. TIMSS & PIRLS International Study Center, Boston College.
- Organisation for Economic Co-operation and Development [OECD]. (2019). *PISA 2018 Assessment and Analytical Framework*. OECD Publishing. <https://doi.org/10.1787/b35f14e5-en>
- Salmerón, L., Strømso, H. I., Kammerer, Y., Stadtler, M., & van den Broek, P. (2018). Comprehension processes in digital reading. In M. Barzillai, J. Thomson, S. Schroeder, & P. van den Broek (Eds.), *Learning to read in a digital world* (pp. 91–120). John Benjamins Publishing.

- Singer, L. M., & Alexander, P. A. (2017). Reading on paper and digitally: What the past decades of empirical research reveal. *Review of Educational Research*, 87(6), 1007–1041. <https://doi.org/10.3102/0034654317723003>
- Stanovich, K. E. (1986). Matthew effects in reading: Some consequences of individual differences in the acquisition of literacy. *Reading Research Quarterly*, 21(4), 360–407. <https://doi.org/10.1598/RRQ.21.4.1>
- Sulkunen, S., Malin, A., & Linnakylä, P. (2010). ICT as a sub-framework in international literacy assessments. *Journal of Writing Research*, 2(2), 115–134.
- Sweller, J. (1988). Cognitive load during problem solving: Effects on learning. *Cognitive Science*, 12(2), 257–285. https://doi.org/10.1207/s15516709cog1202_4
- Sweller, J., Ayres, P., & Kalyuga, S. (2011). *Cognitive load theory*. Springer Science & Business Media.
- van Staden, S., & Bosker, R. (2014). Factors that affect reading achievement: Evidence from South Africa’s PIRLS 2011 results. *South African Journal of Education*, 34(3), 1–9. <https://doi.org/10.15700/201409161058>
- Wästlund, E., Reinikka, H., Norlander, T., & Archer, T. (2005). Effects of VDU and paper presentation on consumption and production of information. *Computers in Human Behavior*, 21(2), 377–388. <https://doi.org/10.1016/j.chb.2004.02.007>
- Wolf, M. (2018). *Reader, come home: The reading brain in a digital world*. Harper.