

RESEARCH ARTICLE

Journal of Global Humanities and Social Sciences
Generative AI in the Humanities: Methods, Applications, and Critical Perspectives
2026, Vol. 1 10-23
DOI: [10.61360/BoniGHSSAH262020660102](https://doi.org/10.61360/BoniGHSSAH262020660102)

From Machine Translation to Human-AI Collaboration: Translator Roles and Responsibilities in Chinese and English-Language Scholarship, 2016–2026



Dandan Wu^{1,*} & Maoli Tian¹

¹*Beijing International Studies University, China.*

Abstract: Generative artificial intelligence is recasting translation as a form of mediated production, yet discussions of what translators are responsible for remain dispersed across technical, pedagogical, professional and ethical literatures. This study maps the development of that discussion through a comparative bibliometric analysis of 447 CNKI records and 303 Web of Science records published between 2016 and 2026. The corpus was assembled from database exports, manually screened, corrected for an Early Access/final-version duplicate, and analysed through publication counts, standardized keyword frequencies, keyword co-occurrence and period-based thematic prevalence. Publication output rose sharply after 2022 in both databases. Chinese-language scholarship gives greater prominence to post-editing, translator subjectivity, human-AI collaboration and translation ethics, whereas the WOS corpus remains more strongly organized around machine translation, neural machine translation, AI, translator training and translation technology. The comparison suggests that generative AI has not displaced the translator from the research field. It has moved the centre of attention from the production of a target text to the supervision of a socio-technical process. The article concludes with a four-dimensional responsibility framework centred on quality, cultural mediation, ethical accountability and professional reconfiguration.

Keywords: generative artificial intelligence, translator agency, translation ethics, post-editing, human-AI collaboration

1. Introduction

The arrival of large language models has altered the vocabulary through which translation is discussed. Earlier research typically asked how machine translation could improve speed, fluency or productivity. More recent work asks what happens when a system can generate a plausible translation, explain a choice, revise its own output and participate in a longer workflow. The technical change matters, but it does not by itself settle the social question. A translation is still addressed to readers, embedded in institutions, shaped by purposes and exposed to consequences. The question is therefore not simply whether a machine can produce a target-language sequence. It is who is expected to judge that sequence, on what

grounds, and with what responsibility when the result is misleading, culturally inappropriate or ethically compromised.

Translation Studies has long had conceptual resources for this question. Translator Studies places the agents of translation, their decision-making and their social conditions alongside the study of translated texts (Chesterman, 2009). Research on machine translation and post-editing has shown that technological assistance changes the distribution of tasks rather than eliminating the need for expertise (Koponen, 2016; Mellinger, 2017). Studies of professional practice have also documented both resistance to and accommodation of machine translation, as well as the anxiety produced by automation and market

Corresponding Author: Dandan Wu

Beijing International Studies University, China

©The Author(s) 2026. Published by BONI FUTURE DIGITAL PUBLISHING CO., LIMITED This is an open access article under the CC BY License (<https://creativecommons.org/licenses/by/4.0/>).

pressure (Cadwell et al., 2018; Vieira, 2020). More recent work extends this debate to machine-translation literacy, posthumanist translation and ethical issues in literary translation (Li, 2023; Ehrensberger-Dow et al., 2023; Lee, 2024).

What remains less clear is how these strands have developed in relation to one another, and whether the pattern is the same in Chinese-language and English-language scholarship. A database comparison is useful because it makes visible not only the growth of a field but also the vocabularies through which different research communities define its central problem. Bibliometric analysis cannot replace close reading of individual studies. It can, however, reveal the scale, timing and semantic organization of a research conversation when the corpus is large enough to exceed the limits of narrative literature review (Aria & Cuccurullo, 2017; Donthu et al., 2021).

This article addresses four questions. RQ1 asks how research output on AI, translation technology and translator roles changed between 2016 and 2026 in CNKI and WOS. RQ2 asks which keywords and keyword pairings organize the two corpora. RQ3 asks how the balance among machine-translation technology, post-editing, translator agency, ethics, competence and human–AI collaboration changed between the 2016–2021 and 2022–2026 periods. RQ4 asks what conception of translator responsibility can be developed from the comparison without reducing responsibility to either technical quality control or a general moral warning about AI.

The article makes three contributions. First, it provides a transparent comparison based on a corrected corpus rather than on a single database. Second, it distinguishes frequency from interpretation. Standardized keyword counts show what is named, while period-based category prevalence shows how the named concepts are distributed across records. Third, it connects the bibliometric findings to a responsibility framework in which the translator is neither a residual human left after automation nor an all-powerful ethical gatekeeper. The translator is treated as a situated mediator who must make and account for decisions across a changing workflow.

2. Conceptual and Analytical Framework

2.1 From machine translation to human–AI collaboration

The period covered by this study contains at least three overlapping technological configurations. Statistical machine translation was followed by neural machine translation, whose attention-based architectures helped make fluent output more widely available (Vaswani et al., 2017). The professional response was not the disappearance of translation work but the growth of post-editing, interactive translation and technology-mediated revision. In this configuration, the system supplied a draft and the translator supplied correction, comparison and quality judgement. Post-editing research consequently treated effort, productivity and quality as related but non-identical dimensions (Koponen, 2016). ISO 18587:2017 formalized requirements for full human post-editing of machine-translation output, thereby recognizing that human intervention remains part of the service rather than an accidental repair.

The second configuration is marked by the integration of machine translation into professional workflows and translator education. Technology competence becomes more than the ability to operate a computer-assisted translation tool. It includes knowing when an automated suggestion is useful, how to interrogate it, how to manage terminology and data, and how to explain the final decision. Massey and Ehrensberger-Dow (2017) and Mellinger (2017) therefore describe machine learning and machine translation as issues for translator education, not only for computer science. Ehrensberger-Dow et al. (2023) sharpen this point by proposing the translator or trainer as an MT literacy consultant.

The third configuration is generative AI. Large language models do not simply retrieve segments or rank alternatives. They produce extended textual responses whose fluency can obscure uncertainty, provenance and uneven performance. Bender et al. (2021) drew attention to the social and environmental costs of large language models, while Floridi and Cowls (2019), UNESCO (2021) and Tabassi (2023) provide broader frameworks for human oversight, transparency, fairness and risk management. In translation, the significance of generative AI lies in the way it expands the space

of possible interventions. The translator may now select among outputs, construct prompts, check sources, compare versions, control terminology, revise cultural references and document a decision trail. The task is less adequately described as correcting a machine sentence than as supervising a chain of mediated choices.

2.2 Translator agency, responsibility and situated judgement

Agency is used here in a modest and operational sense. It refers to the capacity to make consequential choices within institutional, technological and economic constraints. This formulation avoids two opposite assumptions. The first treats the translator as an autonomous author whose decisions are independent of tools, clients and norms. The second treats the translator as a passive operator who merely accepts or rejects a machine output. Translator Studies has emphasized the value of investigating agency through sociology, culture and cognition together (Chesterman, 2009). Research on adaptive expertise similarly presents translators as professionals who reorganize their practices as technologies and communication environments change (Enríquez Ráido, 2016).

Responsibility adds an answerability dimension to agency. A translator can exercise agency without having complete control over a workflow, but a professional role normally carries an expectation that decisions can be justified to clients, readers, institutions or affected communities. Pym's account of translator ethics places mediation between cultures at the centre of the translator's professional situation (Pym, 2012). In the AI context, responsibility therefore concerns more than whether the final text is grammatical. It includes the choice of system, the handling of confidential or copyrighted material, the recognition of uncertainty, the treatment of culturally situated meanings and the disclosure of relevant human or machine intervention.

The responsibility framework developed in this article is consequently relational. It is not a list of personal virtues added to a technical workflow. It links the translator's judgement to the properties of the task, the intended use of the translation and the consequences of error. The framework has four dimensions. Quality responsibility concerns fitness for purpose and the detection of errors that automated metrics may miss. Cultural and

interpretive responsibility concerns the mediation of references, voices and values. Ethical and accountable responsibility concerns data, bias, copyright, confidentiality, transparency and human oversight. Professional responsibility concerns the acquisition and redistribution of expertise as the workflow changes. The dimensions are analytically distinct but practically connected.

2.3 Bibliometric comparison as an analytical method

Bibliometric analysis treats publications and their metadata as traces of scholarly activity. In co-word analysis, words that occur together are interpreted as indicators of conceptual association and the development of problematic networks (Callon et al., 1983). Science mapping adds a visual and comparative dimension to these associations. The method is particularly appropriate when a field contains several partially overlapping vocabularies, because it can show whether terms such as post-editing, translator training, ethics and human–AI collaboration occupy the same research space or remain separated.

The present study follows the methodological logic of bibliometric research described by Zupic and Čater (2015), Aria and Cuccurullo (2017) and Donthu et al. (2021), while limiting the claims to the metadata available in the two databases. The analysis does not infer the quality of a research field from publication volume. Nor does it treat the frequency of a keyword as evidence that all papers using it share a single theoretical position. Instead, publication counts address development, standardized frequencies address naming practices, co-occurrence addresses local association, and category prevalence addresses broad thematic emphasis. This separation is important for a comparison between databases whose indexing rules, languages and citation practices are not identical.

3. Corpus and Methods

3.1 Data sources and search scope

The study uses two bibliographic corpora. The first was assembled from CNKI, with Chinese-language journal records. The second was assembled from Web of Science Core Collection exports, with English-language records indexed as journal articles or related scholarly records in the supplied files. The search window was 2016–2026.

The period was divided at 2021/2022 because the earlier years are dominated by machine translation, neural machine translation and post-editing, whereas the later years contain the visible expansion of ChatGPT, large language models and generative AI. The division is an analytical convention, not a claim that one technology replaced another at a single moment.

The search strategy combined an AI or translation-technology block with a translator, role, responsibility or human–AI block. The CNKI search vocabulary used the Chinese equivalents of artificial intelligence, generative AI, ChatGPT, large language model, machine translation, neural machine translation, intelligent translation and AI translation, combined with the Chinese equivalents

of translator, translator role, translator subjectivity, translation subject, translator responsibility, translation ethics, post-editing and human–AI collaboration. The WOS vocabulary included artificial intelligence, generative AI, ChatGPT, large language model, machine translation, neural machine translation, AI-assisted translation and computer-assisted translation, combined with translator, translator agency, translator role, translator responsibility, translation ethics, accountability, post-editing and human–AI collaboration. The supplied files also include thematic supplement searches for GenAI, translator roles, and post-editing or collaboration. All exports were merged before analysis.

Table 1. Data sources and analytical corpus

Database	Years	Core record fields	De-duplicated candidates	Final records
CNKI	2016–2026	Title, author, source, keywords, abstract, year, fund	869	447
WOS	2016–2026	Title, author, source, keywords, abstract, year, cited references, citations	698 in deduplicated workbook; 707 in screening register	303

Source: supplied CNKI and WOS exports, cleaned data sheets and screening logs. The WOS screening register contains nine additions not present in the 698-row deduplicated workbook. One of those additions was retained in the final corpus.

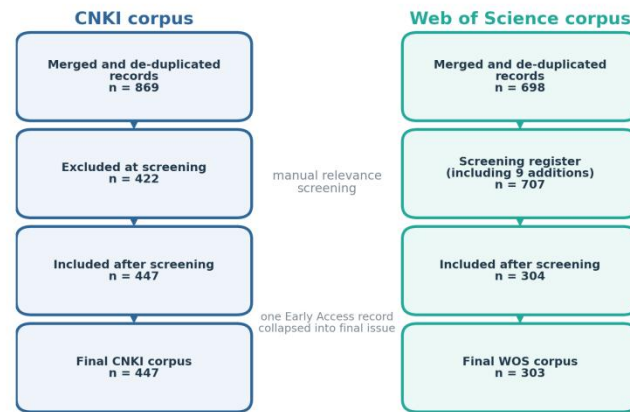
3.2 Corpus construction, screening and duplicate correction

The raw exports were preserved as a separate layer. The analysis proceeded through four stages. Records from the main and supplementary searches were merged within each database. Repeated records were removed by database identifiers where available and by title, author, source and year when needed. The title, abstract and keyword fields were then screened against the inclusion criteria. A record was retained when it addressed translation, translators or translation technology and also engaged at least one of the following themes: translator roles or agency, responsibility or ethics, post-editing, competence or training, or human–AI collaboration. Generic AI or algorithm papers with no meaningful translation connection were excluded.

The reporting flow is presented in a PRISMA-informed manner (Page et al., 2021),

although the study is a bibliometric comparison rather than a systematic review of intervention studies.

The supplied CNKI screening log contained 869 records, of which 447 were marked for inclusion and 422 for exclusion. The supplied WOS screening log contained 707 rows, of which 304 were marked for inclusion. A final audit of the WOS cleaned sheet found the same title twice for Li's 2023 article on ethical issues in literary translation. Record W235 was marked "Article; Early Access" and had no page range, while W241 was marked "Article" with the final issue pages 529–545. The records shared the same author, journal, year and abstract. W235 was therefore removed and W241 was retained. This correction produces a final WOS corpus of 303 distinct records.



Note: The nine WOS additions appear in the supplied screening register but not in the 698-row deduplicated workbook; one of them was retained. Counts follow the corrected final corpus.

Figure 1. Corpus construction and duplicate resolution.

3.3 Measures, keyword normalization and category coding

The unit of analysis is the bibliographic record. Publication output for database d and year y is defined as the number of retained records assigned to that year.

$$N(d, y) = \sum_i I(\text{year}_i = y)$$

For keyword analysis, a keyword was counted once per record after standardization. Thus, machine translation, neural machine translation and related forms were harmonized only where the supplied bilingual mapping identified them as instances of the same concept. The document frequency of keyword k is:

$$F(d, k) = \sum_i I(k \in K_i)$$

Keyword co-occurrence was calculated as the number of records in which two standardized keywords appeared together. Repeated forms within a single record were counted once before pairing.

$$C(d, k, l) = \sum_i I(k \in K_i \text{ and } l \in K_i), \quad k \neq l$$

For the comparative heatmap, standardized keywords were assigned to the conceptual categories recorded in the supplied mapping table. Shared labels such as ChatGPT, Google Translate, post-editing, translation ethics and trust were aligned with the corresponding category in both databases. A category was counted at the article level when at least one of its keywords appeared in the record. The period prevalence of category c is therefore:

$$S(d, p, c) = 100 \times [\sum_i I(\text{year}_i \in p \text{ and } c \in C_i)] / [\sum_i I(\text{year}_i \in p)]$$

The category counts are not mutually exclusive. A record can contribute to machine translation, post-editing and ethics at the same time. The heatmap uses eight substantive categories and excludes the residual “unclassified/other” category from visual comparison. It does not delete those records from the corpus or from keyword-frequency analysis.

Table 2. Corpus screening and correction

Stage	CNKI	WOS
Merged and de-duplicated candidates	869	698
Records in supplied screening register	869	707
Excluded after relevance screening	422	403
Included after relevance screening	447	304
Early Access/final-version correction	0	-1

Stage	CNKI	WOS
Final analysis corpus	447	303

The WOS screening register contains nine records that are not present in the 698-row deduplicated workbook. The final count is based on the corrected record-level corpus, not on an unexamined sum of source files.

3.4 Analytical safeguards

Three safeguards were applied to reduce avoidable error. First, the final record counts were reconciled against the inclusion fields in the screening logs. All 447 CNKI records matched the 447 records marked “Yes”. After the documented WOS duplicate correction, all 303 retained WOS records matched an included record. Second, keyword frequencies were recomputed from the record-level standardized keyword fields and checked against the supplied frequency tables. Third, the two databases were analysed separately before comparison. English labels in the figures are analytical translations of standardized concepts, not additional search terms.

The analysis is descriptive rather than inferential. It reports counts, ratios and document-level prevalence. The phase growth ratio is used only as a compact descriptive indicator:

$$G(d) = N(d, 2022-2026) / N(d, 2016-2021)$$

The CNKI corpus contains 108 records in the first phase and 339 in the second, producing $G(\text{CNKI}) = 3.14$. The WOS corpus contains 90 and 213 records, producing $G(\text{WOS}) = 2.37$. Because 2026 is incomplete, the last year is not interpreted as a final annual total.

4. Results

4.1 Publication trajectories

Both corpora show a marked expansion after 2022, but their trajectories are not identical. CNKI records increase from 16 in 2016 to 29 in 2020, fall to 18 in 2021, and then rise to 32 in 2022, 42 in 2023, 72 in 2024, 109 in 2025 and 84 in the partial 2026 record set. WOS records increase from 10 in 2016 to 24 in 2019, then remain between 16 and 18 from 2020 to 2022 before reaching 39 in 2023, 46 in 2024, 50 in 2025 and 61 in the partial 2026 set.

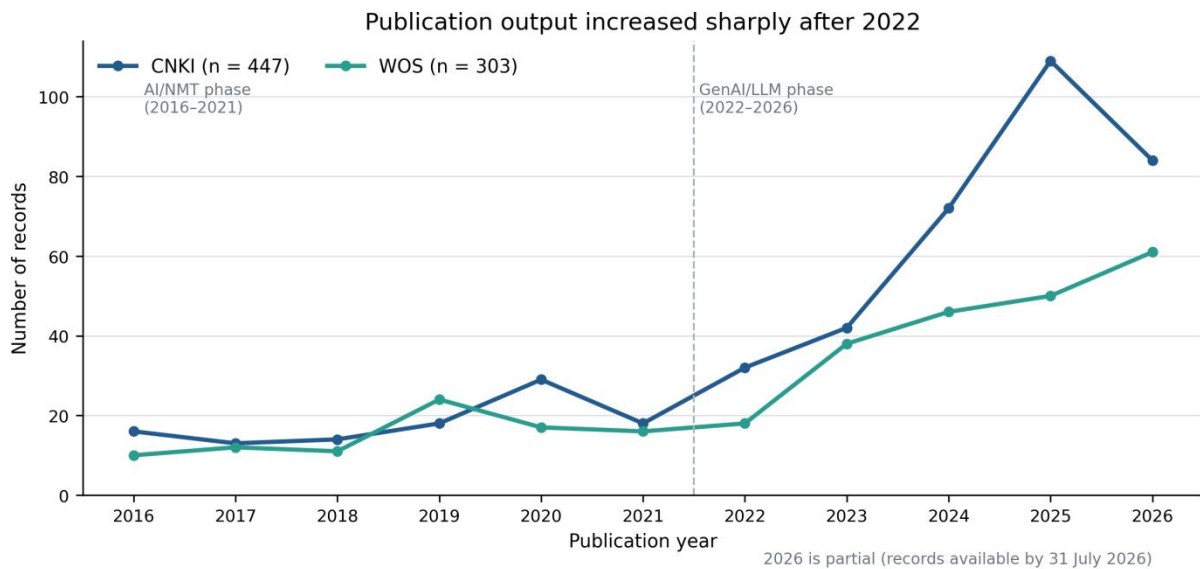


Figure 2. Annual publication trends in the final CNKI and WOS corpora.

The phase comparison captures the scale of the change more clearly than any single year. The second phase contains 75.8% of the CNKI corpus and 70.3% of the WOS corpus. The increase is therefore not limited to a small group of early technical studies. It reflects a wider reorganization

of the research agenda in which generative AI, large language models, translator training and human oversight become increasingly visible. The difference in absolute totals should not be interpreted as a difference in scholarly productivity.

CNKI and WOS index different languages, journals, disciplines and publication practices.

The trajectory also cautions against treating 2022 as a clean boundary. Machine translation and post-editing remain prominent in the later years, while discussion of agency and responsibility predates ChatGPT. The transition is cumulative. Generative AI intensifies an existing debate about human judgement rather than creating that debate from nothing.

4.2 Keyword prominence and co-occurrence

The standardized keyword profiles reveal a shared technical centre and different interpretive

emphases. In CNKI, the six most frequent terms are machine translation (138 records), post-editing (119), artificial intelligence (62), translator subjectivity (51), human–AI collaboration (38) and generative AI (31). In WOS, the leading terms are machine translation (109), AI (37), post-editing (33), neural machine translation (26), translation (23) and large language model (18). Translator training appears in 14 WOS records, while human–AI collaboration appears in 7. The two corpora therefore share machine translation as their most frequent concept, but the surrounding vocabulary is not the same.

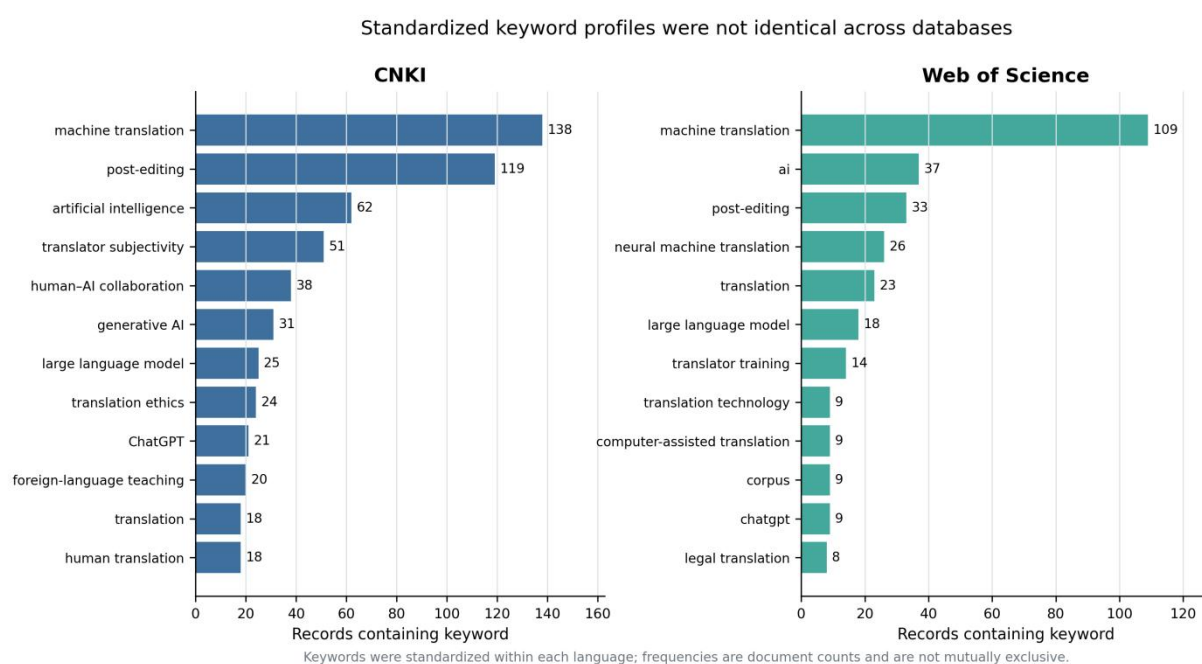


Figure 3. Top standardized keywords in the final corpora.

Table 3. Leading standardized keywords

Rank	CNKI keyword	n	WOS keyword	n
1	machine translation	138	machine translation	109
2	post-editing	119	AI	37
3	artificial intelligence	62	post-editing	33
4	translator subjectivity	51	neural machine translation	26
5	human–AI collaboration	38	translation	23
6	generative AI	31	large language model	18
7	large language model	25	translator training	14

Rank	CNKI keyword	n	WOS keyword	n
8	translation ethics	24	translation technology	9
9	ChatGPT	21	computer-assisted translation	9
10	foreign-language teaching	20	corpus	9

Frequencies are document counts. A record may contain more than one keyword and may therefore appear in several rows.

The most frequent CNKI pair is machine translation and post-editing, which co-occurs in 68 records. Other recurrent pairings include human translation with machine translation (13), artificial intelligence with human–AI collaboration (12), and machine translation with translation quality (8) or artificial intelligence (8). The WOS pair with the highest co-occurrence is also machine translation and post-editing, but it appears in only 24 records. Other WOS pairings include machine translation with neural machine translation (9), audiovisual translation (7), translation (6) and translation technology (6).

This difference is substantive but should not be overstated. Chinese authors more often name the translator’s subjectivity and human–AI collaboration directly in the metadata. WOS records more often retain a technical or professional vocabulary around neural machine translation, translator training and translation technology. The contrast may reflect research emphasis, indexing conventions and keyword habits at the same time. It is best read as a difference in the visibility of concepts rather than as proof that one community lacks the underlying concern.

Table 4. Selected keyword co-occurrences

Database	Most frequent co-occurring pair	n	Second pair	n	Third pair	n
CNKI	machine translation + post-editing	68	human translation + machine translation	13	AI + human–AI collaboration	12
WOS	machine translation + post-editing	24	machine translation + neural machine translation	9	audiovisual translation + machine translation	7

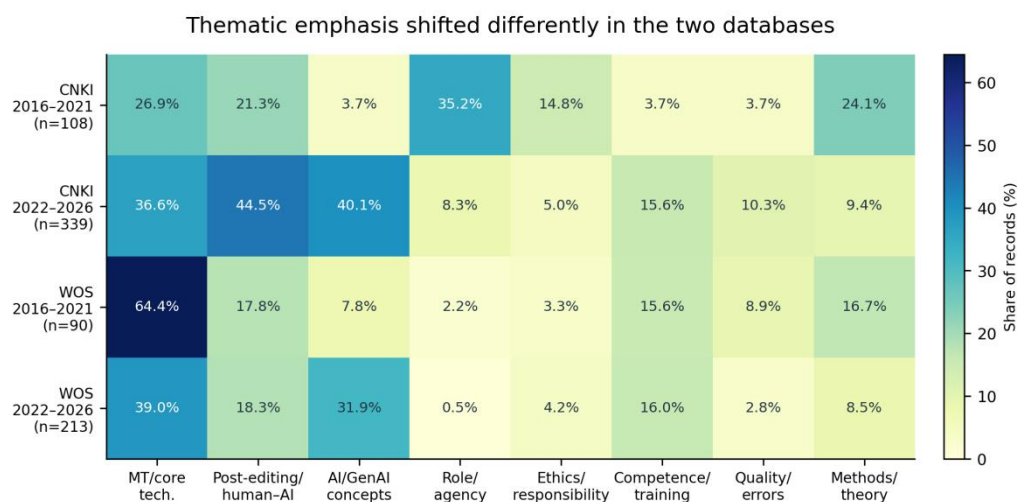
Pairs were counted once per record after within-record de-duplication. The table reports the three highest-frequency pairs in each corpus.

4.3 Thematic change across the two phases

The period-based category analysis provides a more balanced comparison because it uses the number of records in each phase as the denominator. In the CNKI corpus, the share of records containing machine-translation keywords rises from 26.9% in 2016–2021 to 36.6% in 2022–2026. Post-editing and human–AI collaboration rise from 21.3% to 44.5%, while AI and GenAI concepts rise from 3.7% to 40.1%. At the same time, translator role and agency falls from 35.2% to 8.3%, and translation ethics and responsibility falls from 14.8% to 5.0%. The result is not a disappearance of agency or ethics from Chinese scholarship. It indicates that the recent, much larger corpus is more often indexed through

technology and collaboration terms than through the earlier vocabulary of subjectivity and ethics.

The WOS pattern is different. Machine translation is already present in 64.4% of records from 2016–2021 and remains present in 39.0% of records from 2022–2026. AI and GenAI concepts rise from 7.8% to 31.9%. Translator competence and training remains stable at 15.6% in the first phase and 16.0% in the second. Post-editing and human–AI collaboration rises only slightly from 17.8% to 18.3% under the broad category, while explicit translator role and agency is rare in both phases, falling from 2.2% to 0.5%. Translation ethics and responsibility is visible but limited, moving from 3.3% to 4.2



Categories overlap: one article may contribute to several cells. Percentages use the period-specific corpus as denominator.

Figure 4. Thematic prevalence by database and analytical phase.

The comparison suggests two different forms of reorientation. In CNKI, the recent literature moves from a comparatively explicit discussion of the translator as a subject toward a larger research space structured by machine translation, post-editing, GenAI and collaboration. In WOS, the field begins from a stronger technical and NMT base. The later phase adds AI and large-language-model concepts without producing a comparable increase in explicit agency vocabulary. This is a difference in the naming and organization of the problem. It does not mean that WOS authors deny translator judgement. Rather, judgement is more often approached through competence, training, technology adoption, quality or human input.

A further observation concerns the position of ethics. In both corpora, translation ethics and responsibility occupy less space than machine translation and post-editing. The result is striking because AI governance research outside translation has made transparency, accountability, fairness, human oversight and risk management central concepts (Floridi & Cowls, 2019; Jobin et al., 2019; UNESCO, 2021; Tabassi, 2023). The translation literature does engage these questions, particularly in work on literary translation and professional practice, but the metadata show that ethical responsibility has not yet become the main organizing vocabulary of the field.

5. Discussion

5.1 From text production to controlled mediation

The first finding is not simply that publication output increased. It is that the object of inquiry changed as output increased. In the earlier phase, the central problem was how a machine translation system performed and how a translator could work with its output. The dominant relation was therefore system–text–post-editor. In the later phase, the relation becomes system–prompt–output–translator–user, with additional questions about data, context and accountability. The translator’s work includes more forms of selection and orchestration than the term post-editing alone can capture.

This shift helps explain why post-editing remains highly frequent in both databases while human–AI collaboration becomes more visible in CNKI. Post-editing names a recognizable professional task. Collaboration names a broader relationship in which a human may formulate instructions, compare alternatives, supervise multiple outputs and feed decisions back into a workflow. The two concepts should not be treated as synonyms. Post-editing emphasizes the revision of an existing machine output. Collaboration emphasizes the distribution of agency across a process. The latter makes questions of responsibility harder, because a final translation

may be shaped by a sequence of prompts, selections and revisions rather than by one identifiable act.

The field's technical vocabulary is therefore not opposed to its humanistic vocabulary. Machine translation and translation quality provide the material conditions of the problem. Translator subjectivity, agency and ethics provide the language for interpreting what remains irreducibly situated in the task. A useful research agenda should connect the two rather than placing technical evaluation and translator studies in separate compartments.

5.2 Why explicit translator responsibility remains underdeveloped

The relatively low frequency of ethics and responsibility terms does not establish that researchers are unconcerned with ethics. It points to a structural asymmetry in the field. Technical concepts are easy to attach to a method, a dataset or an evaluation metric. Responsibility is distributed across a workflow and often appears under other names, including quality, competence, user acceptance, technology adoption, copyright, confidentiality, bias or human input. Metadata consequently make responsibility look smaller than it is when responsibility is not named as a keyword.

There is also a disciplinary asymmetry. General AI ethics frameworks speak about human dignity, fairness, transparency, accountability and oversight, while translation research often addresses the same issues through situated practices such as terminology management, post-editing, translator

education or literary interpretation. The two vocabularies do not automatically meet. A translator may be responsible for recognizing a culturally inappropriate output even when a general AI framework calls the issue fairness or harm. Conversely, a governance framework may require documentation and oversight without explaining how a translator makes a contextual decision in a particular text.

The bibliometric pattern therefore supports a more specific conception of responsibility. Responsibility in AI-mediated translation should not be reduced to a final human signature placed on a machine-produced text. It should be distributed across identifiable decision points. The translator may not control the training data, the model architecture or the client's commercial deadline, but the translator can often decide whether to use a system, which output to retain, what to verify, how to handle uncertainty and how to communicate limitations. These are professional decisions with ethical consequences.

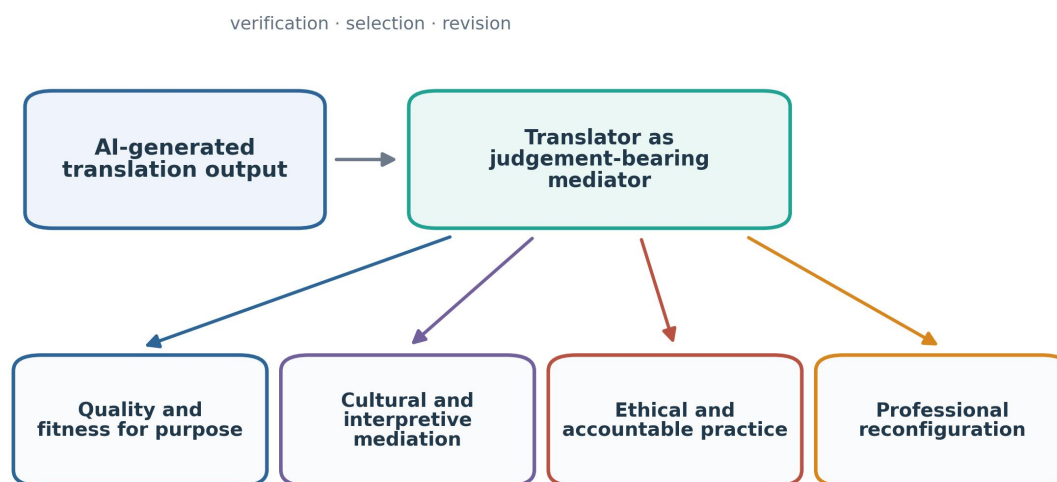
5.3 A four-dimensional framework for translator responsibility

The four dimensions proposed below synthesize the observed keyword and category patterns with established discussions of translator agency, post-editing, professional competence and AI governance. They are not a new competency standard and should not be read as a claim that every study in the corpus addresses all four dimensions. Their purpose is to make the implicit structure of the field more explicit.

Table 5. Four dimensions of translator responsibility in AI-mediated translation

Dimension	Empirical signal in the corpus	Core responsibility	Indicative sources
Quality and fitness for purpose	Machine translation, post-editing, translation quality, error analysis	Judge whether an output is fit for the intended use, identify errors that automated scores miss, and document consequential revisions.	ISO 18587:2017; Koponen, 2016
Cultural and interpretive mediation	Translator subjectivity, literary translation, audiovisual translation, cultural and text-type categories	Preserve, adapt or explain meanings that depend on context, voice, reference and cultural positioning.	Chesterman, 2009; Lee, 2024; Li, 2023
Ethical and accountable practice	Translation ethics, AI ethics, trust, responsibility, copyright and human oversight	Manage data and confidentiality, recognize bias and uncertainty, disclose relevant intervention, and provide an accountable rationale for decisions.	Pym, 2012; Bender et al., 2021; UNESCO, 2021; Tabassi, 2023
Professional reconfiguration	Translator training, competence, technology, human–AI collaboration and automation	Maintain language and domain expertise while acquiring AI literacy, workflow literacy and the ability to supervise tools critically.	Massey & Ehrensberger-Dow, 2017; Mellinger, 2017; Ehrensberger-Dow et al., 2023; Vieira, 2020

The framework is interpretive. The empirical signal identifies the principal metadata categories, while the responsibility statements connect those categories to professional judgement.



The framework is an interpretive synthesis of the corpus patterns, not a claim that every article addresses all four dimensions.

Figure 5. From AI-generated output to a four-dimensional responsibility framework.

The framework changes how translator agency is evaluated. Agency is not measured by the number of words typed without technological assistance. It is visible in the quality of the decisions that structure a mediated workflow. A translator can exercise strong agency while accepting a machine suggestion, provided that acceptance is informed, purposeful and accountable. Conversely, a translator can appear to be the final human author while exercising little agency if the workflow makes unexamined machine output difficult to question.

The framework also clarifies the role of translator education. Training for generative AI cannot be limited to prompt writing or tool operation. It must develop the ability to formulate a translation purpose, assess the conditions under which a model is likely to fail, compare outputs, verify information, preserve cultural and stylistic distinctions, protect sensitive data and explain a final decision. This is consistent with the earlier call for machine-translation competence and MT literacy, but it gives the competence an explicit responsibility dimension (Massey & Ehrensberger-Dow, 2017; Mellinger, 2017; Ehrensberger-Dow et al., 2023).

6. Conclusion

This study compared 447 CNKI records and 303 WOS records on AI, translation technology, translator roles and responsibility between 2016 and 2026. The final corpus was constructed from supplied database exports, screening logs and standardized keyword tables. A documented WOS Early Access/final-version duplicate was removed before analysis. The resulting comparison shows a sharp expansion after 2022 in both databases, with the increase more pronounced in CNKI. It also shows a shared technical centre around machine translation and post-editing, surrounded by different metadata vocabularies. Chinese-language scholarship makes translator subjectivity, human–AI collaboration and translation ethics more visible. The WOS corpus is more strongly organized around machine translation, neural machine translation, AI, translator training and translation technology.

The main implication is conceptual. Generative AI does not make the translator disappear from the research field. It makes the translator's responsibility

more difficult to isolate because responsibility is distributed across systems, prompts, outputs, revisions, institutions and users. Translation research can respond by bringing its humanistic vocabulary into closer dialogue with AI governance and by treating quality, cultural mediation, ethics and professional reconfiguration as connected dimensions of practice.

The findings should be read with three limits in mind. CNKI and WOS are not equivalent windows onto the world, and their records should not be interpreted as a complete representation of Chinese or global research. Keyword analysis is sensitive to indexing and authorial naming, and the category coding used here is a transparent metadata procedure rather than a substitute for full-text interpretation. Finally, 2026 is incomplete. These limits do not weaken the value of the comparison. They define the level at which the results are valid. The study maps the organization and changing visibility of a research conversation. It does not claim to measure the actual quality or ethical performance of every translation workflow.

Future work can extend the corpus beyond the two databases, add full-text analysis, trace citations and institutional collaboration, and examine how responsibility is described in empirical studies of translators rather than only in their metadata. Such work would make it possible to test whether the four dimensions proposed here correspond to observable decision practices. The bibliometric evidence already indicates why that question matters. The field is no longer asking only how machines translate. It is increasingly asking how humans should remain answerable when translation is produced with machines.

Declarations

Competing interests: The authors declare no competing interests.

Ethics approval: Ethics approval was not required because the study analysed bibliographic metadata and did not involve human participants.

Data availability: The study was based on bibliographic exports from CNKI and Web of Science. The cleaned record-level tables, keyword mappings, screening logs and figure inputs are retained with the

research materials. Redistribution of database-derived records remains subject to the terms of the respective databases.

Conflict of Interest

The authors declare that they have no conflicts of interest to this work.

References

- 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>
- 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). Association for Computing Machinery. <https://doi.org/10.1145/3442188.3445922>
- Li, B. (2023). Ethical issues for literary translation in the era of artificial intelligence. *Babel: Revue internationale de la traduction / International Journal of Translation*, 69(4), 529–545. <https://doi.org/10.1075/babel.00334.li>
- Cadwell, P., O'Brien, S., & Teixeira, C. S. C. (2018). Resistance and accommodation: Factors for the (non-) adoption of machine translation among professional translators. *Perspectives: Studies in Translation Theory and Practice*, 26(3), 301–321. <https://doi.org/10.1080/0907676X.2017.1337210>
- Callon, M., Courtial, J.-P., Turner, W. A., & Bauin, 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>
- Chesterman, A. (2009). The name and nature of Translator Studies. *HERMES: Journal of Language and Communication in Business*, 22(42), 13–22. <https://doi.org/10.7146/hjlc.v22i42.96844>
- Donthu, N., Kumar, S., Mukherjee, D., Pandey, N., & Lim, W. M. (2021). How to conduct a bibliometric analysis: An overview and guidelines. *Journal of Business Research*, 133, 285–296. <https://doi.org/10.1016/j.jbusres.2021.04.070>
- Ehrensberger-Dow, M., Delorme Benites, A., & Lehr, C. (2023). A new role for translators and trainers: MT literacy consultants. *The Interpreter and Translator Trainer*, 17(3), 393–411. <https://doi.org/10.1080/1750399X.2023.2237328>
- Enríquez Raido, V. (2016). Translators as adaptive experts in a flat world: From globalization 1.0 to globalization 4.0? *International Journal of Communication*, 10, 970–988. <https://ijoc.org/index.php/ijoc/article/view/3504>
- Floridi, L., & Cowls, J. (2019). A unified framework of five principles for AI in society. *Harvard Data Science Review*, 1(1). <https://doi.org/10.1162/99608f92.8cd550d1>
- International Organization for Standardization. (2017). *ISO 18587:2017: Translation services—Post-editing of machine translation output—Requirements*. ISO.
- Jobin, A., Ienca, M., & Vayena, E. (2019). The global landscape of AI ethics guidelines. *Nature Machine Intelligence*, 1, 389–399. <https://doi.org/10.1038/s42256-019-0088-2>
- Koponen, M. (2016). Is machine translation post-editing worth the effort? A survey of research into post-editing and effort. *The Journal of Specialised Translation*, (25), 131–148. <https://doi.org/10.26034/cm.jostrans.2016.303>
- Lee, T. K. (2024). Artificial intelligence and posthumanist translation: ChatGPT versus the translator. *Applied Linguistics Review*, 15(6), 2351–2372. <https://doi.org/10.1515/applirev-2023-0122>
- Massey, G., & Ehrensberger-Dow, M. (2017). Machine learning: Implications for translator education. *Lebende Sprachen*, 62(2), 300–312. <https://doi.org/10.1515/les-2017-0021>
- Mellinger, C. D. (2017). Translators and machine translation: Knowledge and skills gaps in translator pedagogy. *The Interpreter and Translator Trainer*, 11(4), 280–293. <https://doi.org/10.1080/1750399X.2017.1359760>
- Page, M. J., McKenzie, J. E., Bossuyt, P. M., Boutron, I., Hoffmann, T. C., Mulrow, C. D., Shamseer, L., Tetzlaff, J. M., Akl, E. A., Brennan, S. E., Chou, R., Glanville, J., Grimshaw, J. M., Hróbjartsson, A., Lalu, M. M., Li, T., Loder, E. W., Mayo-Wilson, E., McDonald, S., ... Moher, D. (2021). The PRISMA 2020 statement: An updated guideline for reporting systematic reviews. *BMJ*, 372, Article n71. <https://doi.org/10.1136/bmj.n71>
- Pym, A. (2012). *On translator ethics: Principles for mediation between cultures* (Rev. ed.; H. Walker, Trans.). John Benjamins.
- Tabassi, E. (2023). *Artificial intelligence risk management framework* (AI RMF 1.0). National Institute of Standards and Technology. <https://doi.org/10.6028/NIST.AI.100-1>
- UNESCO. (2021). *Recommendation on the ethics of artificial intelligence*. UNESCO.
- van Eck, N. J., & Waltman, L. (2010). Software survey: VOSviewer, a computer program for bibliometric mapping. *Scientometrics*, 84(2), 523–538. <https://doi.org/10.1007/s11192-009-0146-3>
- Vaswani, A., Shazeer, N., Parmar, N., Uszkoreit, J., Jones, L., Gomez, A. N., Kaiser, Ł., & Polosukhin,

- I. (2017). Attention is all you need. In *Advances in Neural Information Processing Systems 30* (pp. 5998–6008). Curran Associates, Inc.
- Vieira, L. N. (2020). Automation anxiety and translators. *Translation Studies*, 13(1), 1–21. <https://doi.org/10.1080/14781700.2018.1543613>
- Waltman, L., van Eck, N. J., & Noyons, E. C. M. (2010). A unified approach to mapping and clustering of bibliometric networks. *Journal of Informetrics*, 4(4), 629–635. <https://doi.org/10.1016/j.joi.2010.07.002>
- Zupic, I., & Čater, T. (2015). Bibliometric methods in management and organization. *Organizational Research Methods*, 18(3), 429–472. <https://doi.org/10.1177/1094428114562629>

Appendix A. Search and processing record

The project files preserve the following processing layers: raw CNKI HTML exports and raw WOS plain-text exports; database-specific de-duplicated workbooks; screening logs; cleaned data sheets; bilingual keyword mappings; and figure-input summaries. The final analysis used the cleaned CNKI data after reconciliation with the CNKI screening log and the cleaned WOS data after removal of the W235 Early Access duplicate. The figure inputs reproduce all annual counts, top-keyword counts and heatmap percentages reported in the article.

How to Cite: Wu, D. & Tian, M. (2026). From Machine Translation to Human – AI Collaboration: Translator Roles and Responsibilities in Chinese and English-Language Scholarship, 2016 – 2026. *Journal of Global Humanities and Social Sciences*, 1(Special Issue), 10-23 <https://doi.org/10.61360/BoniGHSSAH262020660102>