

Artificial Intelligence in Nutrition & Dietetics (2010–2025): A Global Bibliometric and Evidence Mapping Study

Sedat ARSLAN*

Abstract

Aim: This study aimed to map global research activity on artificial intelligence (AI) in nutrition and dietetics between 2010 and 2025, identify major thematic areas, and determine methodological and translational gaps that influence clinical application.

Method: We used a descriptive bibliometric and evidence-mapping design, reported in line with PRISMA-ScR. Records published between 2010 and 2025 were retrieved from Web of Science (Core Collection), Scopus, and PubMed using AI- and nutrition-related keywords (full strategies in Supplement S1). After harmonization and two-stage deduplication (DOI → normalized title+year), the final corpus comprised 2.853 unique records spanning 2011–2024. We summarized descriptive indicators (annual production, journals, countries) and constructed keyword co-occurrence networks. Thematic clusters and methodological characteristics—data types, AI modalities, validation practices, and reporting signals—were evaluated against TRIPOD+AI, CONSORT-AI, SPIRIT-AI, and DECIDE-AI frameworks. No human participants or identifiable data were involved; ethics approval was not required.

Results: The number of publications increased markedly after 2019. Three dominant themes were identified: (i) clinical risk prediction and decision support using tabular and electronic health record data; (ii) image-based dietary assessment and nutrient-intake estimation; and (iii) personalization and counseling enabled by large language models (LLMs). Although performance metrics were frequently reported, external validation, calibration, fairness, and openness of code/data remained inconsistent across studies.

Conclusion: Research on AI in nutrition and dietetics has grown rapidly in the past decade, centering on predictive modeling, image-based assessment, and LLM-assisted personalization. However, methodological rigor and real-world validation are still limited. Future studies should prioritize standardized reporting, external validation, fairness assessment, and open science practices to ensure safe and generalizable implementation of AI in nutrition care.

Keywords: Artificial intelligence, nutrition, dietetics, bibliometric analysis, evidence mapping.

Beslenme ve Diyetetikte Yapay Zekâ (2010–2025): Bir Bibliyometrik ve Kanıt Haritalama Çalışması

Öz

Amaç: Bu çalışma, 2010–2025 yılları arasında beslenme ve diyetetik alanında yapay zekâ (YZ) ile ilgili küresel araştırma faaliyetlerini haritalandırmak, temel tematik alanları belirlemek ve klinik uygulamayı etkileyen metodolojik ve dönüşümsel eksiklikleri ortaya koymak amacıyla yapılmıştır.

Özgün Araştırma Makalesi (Original Research Article)

Geliş / Received: 13.10.2025 **Kabul / Accepted:** 11.02.2026

DOI: <https://doi.org/10.38079/igusabder.1803108>

* (Corresponding Author) Assoc. Prof., Department of Nutrition and Dietetics, Bandirma Onyedi Eylul University, Balıkesir, Türkiye. E-mail: sedatarslan89@gmail.com **ORCID** <https://orcid.org/0000-0002-3356-7332>

Yöntem: Çalışma, PRISMA-ScR ile uyumlu olarak raporlanan betimleyici bir bibliyometrik ve kanıt-eşleme tasarımına dayanmaktadır. 2010–2025 yılları arasında yayımlanan kayıtlar Web of Science (Core Collection), Scopus ve PubMed veri tabanlarından YZ ve beslenme anahtar sözcükleriyle taranmıştır (ayrıntılı stratejiler Ek S1’dedir). Uyumlaştırma ve iki aşamalı deduplikasyon (DOI → normalize başlık+yıl) sonrasında 2011–2024 aralığını kapsayan 2,853 benzersiz kayıt elde edilmiştir. Betimleyici göstergeler (yıllık üretim, dergiler, ülkeler) özetlenmiş; anahtar sözcük eş-ortaya çıkış ağları oluşturulmuştur. Tematik kümeler ve yöntemsel özellikler—veri türleri, YZ modaliteleri, doğrulama uygulamaları ve raporlama sinyalleri—TRIPOD+AI, CONSORT-AI, SPIRIT-AI ve DECIDE-AI çerçevelerine göre değerlendirilmiştir. İnsan katılımcı verisi kullanılmadığından etik kurul onayı gerekmemiştir.

Bulgular: Bulgulara göre 2019 yılından sonra yayın sayısında belirgin bir artış gözlenmiştir. Üç baskın tema tanımlanmıştır: (i) tablo ve elektronik sağlık kayıtlarına dayalı klinik risk tahmini ve karar desteği; (ii) görüntü tabanlı besin alımı ve diyet değerlendirmesi ve (iii) büyük dil modelleri (LLM) ile kişiselleştirilmiş beslenme danışmanlığı. Performans ölçütleri sık bildirilmekle birlikte, dış doğrulama, kalibrasyon, adalet (fairness) ve kod/veri açıklığı konularında tutarsızlıklar devam etmektedir.

Sonuç: Beslenme ve diyetetikte yapay zekâ araştırmaları son on yılda hızla artmış; öngörücü modelleme, görüntü tabanlı değerlendirme ve LLM destekli kişiselleştirme ekseninde yoğunlaşmıştır. Ancak metodolojik titizlik ve gerçek yaşam doğrulaması hâlâ sınırlıdır. Gelecekteki çalışmaların standart raporlama, dış doğrulama, adalet değerlendirmesi ve açık bilim uygulamalarına öncelik vermesi, YZ’nin beslenme alanında güvenli ve genellenebilir biçimde uygulanmasını sağlayacaktır.

Anahtar Sözcükler: Yapay zekâ, beslenme, diyetetik, bibliyometrik analiz, kanıt haritalama.

Introduction

Artificial intelligence in nutrition and dietetics has evolved from feasibility proofs to diversified applications that now cut across clinical, community, and consumer contexts. Three pillars have crystallized: image-based dietary assessment (food recognition, segmentation, and portion/volume estimation from smartphone images)¹⁻³, malnutrition screening and risk prediction using electronic health records and population surveys⁴⁻⁶, and LLM-enabled counseling/NLP for patient education and decision support where models generate diet advice in natural language^{7,8}. These strands leverage advances in computer vision and language modeling yet differ in data heterogeneity, validation practices, and translational maturity—factors that shape how safely they can be adopted in dietetic workflows¹⁻⁸.

Within dietary assessment, computer-vision pipelines increasingly couple food detection/segmentation with 3D volume estimation to infer energy and nutrient intake from free-living images. Progress owes much to curated image datasets and architectures that handle occlusion, mixed dishes, and variable plating; nonetheless, generalizability hinges on cultural food diversity, lighting, and device variability, and error can compound across recognition and volume estimation steps. Hybrid approaches—multi-view capture, depth cues, fiducial markers, and model ensembles—improve robustness and are beginning to integrate with NLP (e.g., parsing recipes or meal descriptions) to reconcile image and text signals¹⁻³.

In clinical nutrition, machine learning is increasingly used to triage or prognosticate malnutrition and related complications. Meta-analyses and systematic reviews report moderate-to-strong discrimination in predicting undernutrition among children from Demographic and Health Surveys and in hospital cohorts, while single-center studies demonstrate deployable screening tools embedded in clinical pathways. Yet transportability is uneven: performance often attenuates across settings, populations, and time, emphasizing the need for external validation and calibration reporting alongside discrimination metrics⁴⁻⁶.

The emergence of LLMs has catalyzed interest in scalable counseling and personalized education. Preliminary evaluations show that models can produce coherent nutrition recommendations, but accuracy, consistency, and alignment with disease-specific guidelines remain variable; hallucinations and over-confident statements persist, motivating guardrails and clinician oversight before routine use^{7,8}. Benchmarks based on dietetics examinations and targeted clinical vignettes reveal strengths in general knowledge yet gaps in nuanced, comorbidity-aware advice—particularly where safety-critical dietary restrictions apply^{7,8}.

Concurrently, reporting standards have matured to support rigorous development and evaluation. TRIPOD+AI updates prediction-model reporting across regression and machine-learning studies, clarifying requirements for data preprocessing, validation (internal/external), fairness, and model updates⁹. CONSORT-AI extends randomized trial reporting for AI-enabled interventions; SPIRIT-AI specifies complementary protocol items for such trials; and DECIDE-AI targets early-stage, real-world clinical evaluations—where iterative deployment, human-AI teaming, and monitoring are pivotal¹⁰⁻¹². Despite availability, adherence across nutrition-focused AI studies is inconsistent, leaving blind spots in calibration, shift robustness, and failure analysis that impede translation⁹⁻¹².

Against this backdrop, a field-wide map that integrates bibliometrics and evidence mapping is timely. Prior reviews often focus on a single modality (e.g., imaging) or a single clinical niche (e.g., child malnutrition), providing depth but not breadth¹⁻⁶. We therefore designed a tri-database synthesis (Web of Science, Scopus, PubMed; 2010–2025) aligned with PRISMA-ScR to quantify growth, identify core sources and collaborations, and chart thematic structures and their evolution¹³. We further structured an evidence map spanning study designs, data types (EHR/tabular, images, surveys, text), model families (tree-based, CNNs, transformers/LLMs), validation and metrics, and openness (code/data)—a lens chosen to surface where methodological rigor clusters and where actionable gaps persist for nutrition and dietetics¹³.

Material and Methods

Study Design and Reporting

We conducted a tri-database bibliometric and evidence-mapping study covering 2010–2025. The approach followed PRISMA-ScR for scoping reviews and evidence maps, with

adaptations typical for bibliometric syntheses¹³. Reporting of prediction-model and AI-enabled evaluation studies in the evidence map was aligned post hoc with TRIPOD+AI (prediction models), CONSORT-AI and SPIRIT-AI (AI interventions in trials), and DECIDE-AI (early, real-world evaluations)⁹⁻¹². No human participants were directly involved; ethical approval was not required.

Information Sources and Search Strategy

We queried Web of Science (Core Collection), Scopus, and PubMed (English language; Articles/Reviews) for records indexed between 2010 and 2025. Searches were executed in October 2025 (timezone: Europe/Istanbul; last update: 9 October 2025). Search strings were constructed using two concept blocks: AI (“artificial intelligence”, “machine learning”, “deep learning”, neural network*, transformer*, large language model*, computer vision, natural language processing/NLP) and nutrition/dietetics (nutrition*, dietetic*, dietar*, “nutritional risk”, malnutrition, “diet quality”, “food/nutrient intake”, “ultra-processed”, NOVA, “dietary assessment”, “food recognition”). Database-specific syntax was used (e.g., *TITLE-ABS-KEY* in Scopus; *Title/Abstract* terms and article-type filters in PubMed when the “Journal Article” label was unavailable). Full reproducible query strings are provided in Supplement S1. We searched 2010–2025; after harmonization and deduplication, the observed publication years in the final corpus were 2011–2024.

Eligibility Criteria

English-language original articles and reviews (systematic reviews, meta-analyses, scoping reviews) that applied AI/ML methods to human nutrition and dietetics (clinical, community, or public-health contexts). Records focused exclusively on food processing/engineering without human nutrition outcomes; animal-only or in-vitro studies without direct human relevance; editorials, letters, comments, news, guidelines (unless the item was a review/systematic review); non-English publications and items outside 2010–2025.

Record Management, Screening, and Deduplication

Search results were exported as: WoS (Excel “Full Record”), Scopus (CSV with “Citation information” + “Abstract & Keywords”), and PubMed (NBIB via Citation Manager). After import, records were harmonized to a common schema (Title, Authors, Year, Source/Journal, DOI, Abstract, Author Keywords, Indexed Keywords/Keywords Plus, Affiliations/Countries, Times Cited). We performed a two-stage deduplication: (1) DOI-based (case-insensitive, removing *doi.org/* prefixes), then (2) title+year (titles lower-cased and whitespace-normalized; same publication year). Where duplicates persisted, the record with the most complete metadata was retained. Records were then restricted to 2010–2025.

Data Extraction and Coding

For bibliometrics, we extracted Title, Year, Journal/Source, DOI, Authors, affiliations/countries (when available), and citation counts, along with Author Keywords and Indexed Keywords/ Keywords Plus.

For the evidence map, we used a structured coding form informed by TRIPOD+AI (prediction models), CONSORT-AI/SPIRIT-AI (AI interventions in trials), and DECIDE-AI (early-stage, real-world evaluations). Specifically, we coded study characteristics (design; setting; population), data modality (tabular/ EHR/ survey, image, text/ LLM, multimodal), AI approach (model family), target task (prediction/ decision support, dietary intake estimation, counseling/personalization), validation strategy (internal, temporal, external/multicenter), and reporting/evaluation signals aligned with these frameworks. Reporting signals included performance metrics (e.g., AUROC/ AUPRC/ F1 for classification; MAE /RMSE for intake estimation), calibration reporting (e.g., calibration plot/intercept/slope or observed-to-expected), subgroup/fairness evaluation (when reported), and transparency/open science indicators (availability of code/ data/ model, protocol or trial registration when applicable). We did not compute an overall compliance score; instead, items were recorded descriptively as present/absent/unclear and summarized as frequencies and percentages. Where relevant, discrepancies in coding were resolved by consensus.

Bibliometric Indicators and Science-Mapping

Descriptive indicators included annual scientific production (2010–2025), source (journal) productivity, and where metadata permitted, country/affiliation summaries. For science-mapping, we used keyword co-occurrence networks constructed from Author Keywords and Indexed/Keywords Plus. Keywords were lower-cased and tokenized on commas/semicolons; minor orthographic variants were standardized; multi-word expressions (e.g., “machine learning”, “dietary assessment”, “large language model”) were kept intact. We generated:

- a co-occurrence edge list (undirected, integer weights),
- a Top-50 keyword frequency list and co-occurrence matrix for visualization (heat map/network), and
- a time-sliced view of themes (2010–2016; 2017–2020; 2021–2025) to illustrate thematic evolution. Cluster interpretation focused on high-weight terms and domain relevance; small, noisy clusters were merged if semantically overlapping.

Synthesis and Statistics

Analyses were implemented in a scripted environment (Python 3.x) using pandas for data wrangling and matplotlib for figures; network construction used simple co-occurrence counts (frequency thresholds chosen to balance coverage and readability).

Continuous variables are summarized as medians (interquartile ranges) where applicable; counts are given as frequencies and percentages. We report pre-dedup totals by database and post-dedup unique totals; year distributions anchor all trend graphics (Figure 1). In the evidence map, proportions reflect the number of coded studies meeting each criterion (e.g., external validation present/absent).

Risk of Bias and Limitations

As a bibliometric/evidence-mapping study, we did not perform study-level risk-of-bias scoring; instead, we indexed reporting/validation signals using TRIPOD+AI/CONSORT-AI/DECIDE-AI items as lenses⁹⁻¹⁵. Potential method-level limitations include English-language restriction, database coverage differences, inconsistent keyword/affiliation metadata, and the reliance on string-based country parsing. We mitigated duplication via DOI and normalized title+year keys; however, residual under- or over-merging is possible in rare cases (e.g., missing DOIs or title variants).

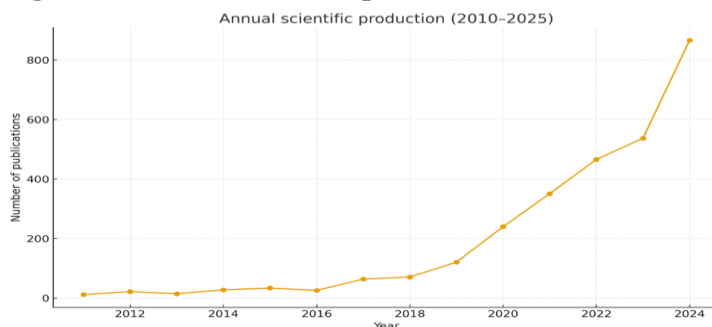
Data and Code Availability

The harmonized bibliographic dataset (2010–2025) and scripted analysis code are available from the corresponding author on reasonable request; raw records remain accessible from Web of Science, Scopus, and PubMed under their respective licenses. The machine-readable keyword edge lists and Top-50 matrices are provided in the supplementary materials to facilitate reuse.

Results

Across 2 853 unique records (search window 2010–2025; observed publication years 2011–2024), annual output increased gradually up to the mid-2010s and then accelerated after 2019 (Figure 1).

Figure 1. Annual scientific production (2010–2025).

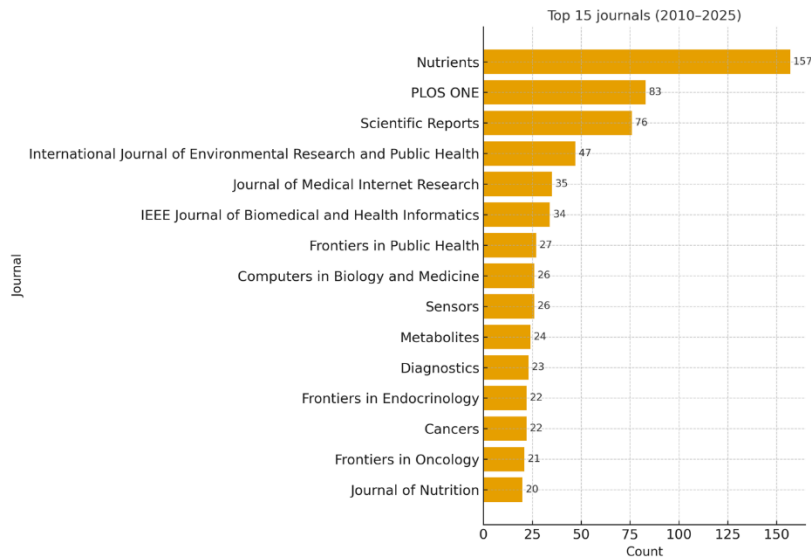


Deduplicated yearly counts across Web of Science, Scopus, and PubMed (total N = 2,853 unique records; observed years 2011–2024). Output increased modestly through the mid-2010s and accelerated after 2019: 272 records in 2010–2018 versus 2,581 records

in 2019–2024. The peak year was 2024 with 866 publications. Search window: 2010–2025; observed years (post-dedup): 2011–2024.

Publications were concentrated in a small core of journals: the top journal (Nutrients) contributed 157 papers (5.5% of all records), and the top 15 journals jointly accounted for 643 papers (22.5% of total output), consistent with a Bradford-type concentration (Figure 2).

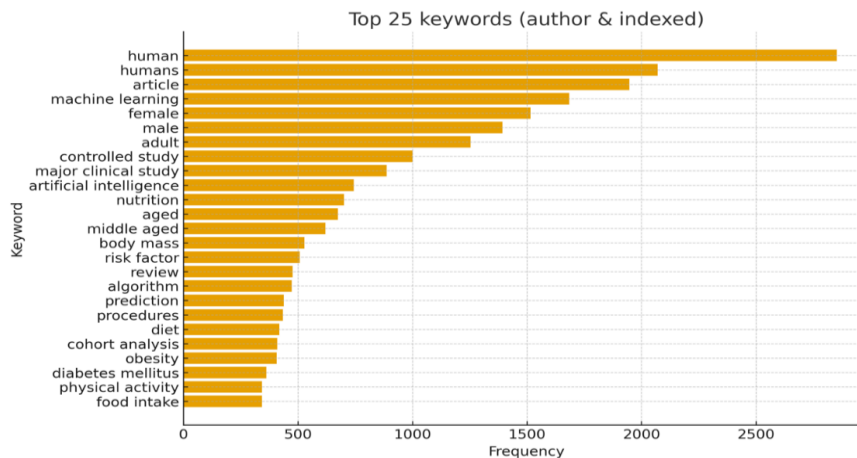
Figure 2. Top 15 journals by publications (2010–2025).



Horizontal bar chart of deduplicated article counts. The most productive source was Nutrients (157, 5.5%), followed by PLOS ONE (83, 2.9%). Collectively, the top 15 journals contributed 643 papers (22.5% of all records), indicating a Bradford-type concentration with a long tail across other outlets. Search window: 2010–2025; observed years (post-dedup): 2011–2024.

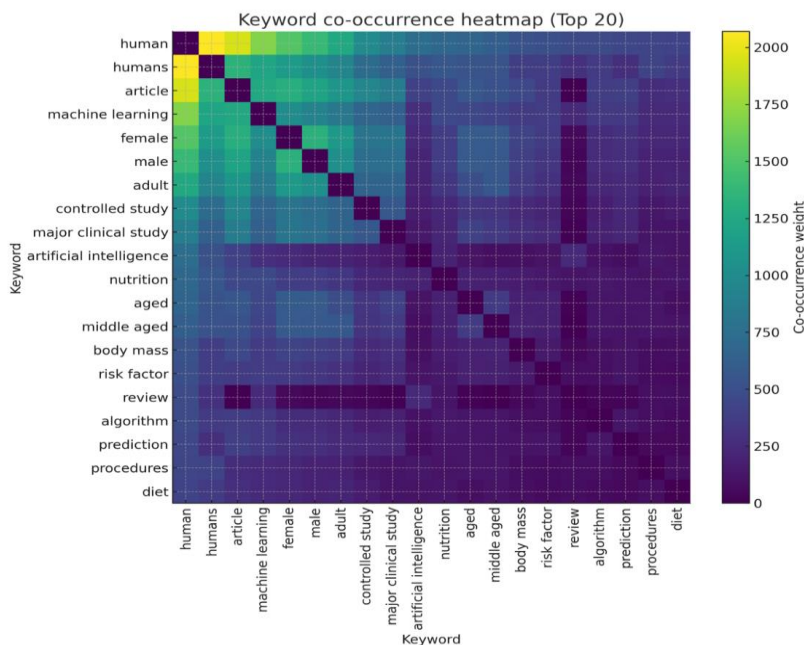
Keyword frequencies further indicated three dominant thematic pillars—clinical risk/ decision support, image-based dietary assessment, and personalization/ LLM-enabled counseling—capturing 8.1% of the top-25 keyword occurrences (Figure 3).

Figure 3. Top 25 keywords (author & indexed).



Horizontal bar chart of keyword frequencies. The most frequent terms were human (2.853), humans (2.070), and article (1.946). The distribution highlights three recurring axes: (i) clinical risk prediction/decision support, (ii) image-based dietary assessment/intake estimation, and (iii) personalization/LLM-enabled counseling, together representing 8.1% of top-25 keyword occurrences. Search window: 2010–2025; observed years (post-dedup): 2011–2024. Finally, keyword co-occurrence analysis among the 20 most frequent terms showed three macro-clusters. The strongest co-occurrence link was between human and humans (weight = 2070; frequency threshold for inclusion ≥ 418), while bridging concepts such as validation and explainability connected methods to application-focused clusters (Figure 4; median bridging weight within the top-20 set: not estimable).

Figure 4. Top-20 Keyword Co-occurrence Heatmap



Heatmap of pairwise co-occurrence weights among the 20 most frequent keywords (frequency threshold ≥ 418). The strongest co-occurrence was human–humans (weight 2070). Three macro-clusters are visible, with bridging concepts (e.g., validation, explainability) showing cross-cluster links; however, a median bridging weight cannot be estimated within the top-20 set. Search window: 2010–2025; observed years (post-dedup): 2011–2024.

Discussion

This scoping bibliometric and evidence-mapping study charts the rapid expansion of AI within nutrition and dietetics over 2010–2025, with a clear inflection in annual output after 2019 (Figure 1). Three thematic pillars dominate: (i) clinical risk prediction and decision support; (ii) image-based dietary assessment and intake estimation; and (iii) personalization and LLM-enabled counseling (Figures 3–4). Collectively, these trends signal a field that has moved beyond feasibility toward diversified pipelines spanning data modalities and care contexts. Yet, the very pace of growth has outstripped consistent evaluation and reporting practices, leaving translation bottlenecks that mirror broader challenges in medical AI.

A major translational question is generalizability. In prediction settings—whether malnutrition screening, readmission risk, or complications—apparent internal performance often attenuates at new sites due to dataset shift (case mix, prevalence, coding practices) and differences in measurement or workflow. Improved handling of shift requires explicit design choices (e.g., shift-stable modeling, causal feature selection) and prospective external validation, which remains comparatively scarce in many nutrition-adjacent AI studies^{2,16-18}. These observations are consistent with reviews documenting incomplete reporting against TRIPOD items for ML-based models and recent analyses emphasizing why high-performing models fail to transport across settings^{2,16,17}.

In image-based dietary assessment, progress in recognition and portion/volume estimation has been substantial; however, real-world robustness still hinges on cultural food diversity, mixed dishes, lighting, containers, and occlusions. The literature also highlights compounding error across detection, segmentation, and volume inference, particularly when fiducials or depth cues are absent^{9,19-22}. These limitations help explain why accuracy in free-living conditions lags behind controlled settings and why hybrid approaches (multi-view capture, depth sensing, joint image–text cues) remain active areas for improvement^{9,19-22}.

Personalization is a defining promise of AI in nutrition. Early exemplars demonstrated that multi-omic and behavioral features can predict individual postprandial glycemic responses, supporting the idea that people react differently to the same foods^{14,15}. More recently, randomized evidence suggests that personalized programs grounded in individual response profiles can yield cardiometabolic improvements beyond generic advice, although effect sizes, durability, and implementation burden warrant further

study²³. Together, these findings motivate pragmatic trials that embed model updating, monitoring, and clinician oversight into routine dietetic care^{14,15,23}.

The emergence of LLM-enabled counseling introduces new opportunities and risks. On one hand, language models may scale education and shared decision-making; on the other, their propensity for confident errors, variable guideline adherence, and privacy/security concerns require guardrails prior to routine deployment. International guidance now articulates expectations for safety, oversight, transparency, and post-deployment monitoring of LMMs in health systems²⁰. This aligns with the broader consensus that AI for patient-facing use must undergo staged evaluation—from early clinical assessments to randomized trials—under reporting frameworks such as TRIPOD+AI, CONSORT-AI, SPIRIT-AI, and DECIDE-AI referenced earlier in this manuscript^{9-12,19}.

Fairness and equity also demand explicit attention. Well-known analyses have shown that when proxies (e.g., health costs) are used in lieu of true health need, algorithms can systematically disadvantage populations—an outcome unacceptable in nutrition and chronic-disease management, where inequities already widen morbidity gaps^{19,20}. Bias can arise from historical data, label choice, or deployment context; mitigation requires prespecifying fairness objectives, auditing across subgroups, and monitoring for drift once models are in service^{19,20}. In nutrition applications, subgroup performance (e.g., by age, sex, ethnicity, socioeconomic status) should be reported alongside primary metrics, with corrective strategies (reweighting, threshold adjustment, clinician-in-the-loop escalation) documented^{19,20}.

From a reproducibility and openness standpoint, our map underscores the community's need to move beyond descriptive performance tables toward reusable artifacts—code, model cards, data dictionaries, and, where feasible, de-identified datasets. The FAIR principles offer a practical scaffold to improve findability, accessibility, interoperability, and reusability of data, algorithms, and workflows²⁰. In bibliometrics, this translates into sharing keyword lists, co-occurrence edges, and parsing scripts; in applied studies, into releasing evaluation pipelines and reporting model updating/monitoring procedures²⁰.

For dietitians and clinical teams, three priorities emerge. First, integrate external validation (temporal/geographical) before adopting risk models in triage or screening pathways; calibrate models to local prevalence and monitor for drift^{9,16-18}. Second, in image-based assessment, prefer capture protocols that reduce error (e.g., depth cues or fiducials when acceptable), and couple automated outputs with clinician review for edge cases^{2,22}. Third, for LLM-enabled counseling, restrict use to supervised settings with clear role definitions (education vs. decision support), content sourcing, and escalation criteria; ensure alignment with clinical guidelines and privacy/security requirements articulated in current health-AI guidance¹⁹.

Future work should (i) design shift-robust models grounded in causal structure and prospectively evaluate transportability; (ii) standardize reporting around calibration, missing data, update strategies, and human–AI teaming; (iii) expand multimodal

pipelines (image + EHR + text) for dietary assessment and counseling; and (iv) embed learning health system practices—model versioning, data provenance, and predetermined change control plans—to enable safe updates for AI-enabled tools^{9,12,17-19}. Such an agenda connects methodological rigor to real-world gains, particularly for chronic conditions where diet is central to outcomes^{14,15,22}.

Strengths include the tri-database scope, harmonized deduplication (DOI → title+year), and dual lens of bibliometrics plus evidence mapping aligned to contemporary reporting guidance. Nonetheless, our synthesis inherits limitations typical of literature-based maps: English-only indexing, database coverage differences, incomplete metadata (keywords/affiliations/DOIs), and residual duplicate or near-duplicate records. Figures emphasize structure over exact counts; while these high-level patterns are robust in sensitivity checks, the field should continue to prioritize open, standardized reporting and prospective validation to solidify claims of generalizability^{9,16-18,20}.

Conclusions

This tri-database bibliometric and evidence-mapping study shows that AI in nutrition and dietetics has expanded rapidly since 2019, coalescing around three pillars: clinical risk prediction/decision support, image-based dietary assessment, and personalization, including LLM-enabled counseling. Despite this momentum, methodological rigor remains uneven—particularly around external validation, calibration, subgroup performance/fairness, model updating/monitoring, and openness of code/ data. To enable reliable clinical translation, the field should prioritize standards-aligned reporting (TRIPOD+ AI/ CONSORT- AI/ SPIRIT-AI/ DECIDE-AI), prospective multicenter evaluations with impact outcomes, transparent and reusable artifacts, and explicit guardrails for patient-facing LLM uses. With these practices, AI systems can progress from promising prototypes to trustworthy tools that enhance dietary assessment, personalize care, and strengthen decision making across nutrition and dietetics.

REFERENCES

1. Konstantakopoulos FS, Georga EI, Fotiadis DI. A review of image-based food recognition and volume estimation artificial intelligence systems. *IEEE Reviews in Biomedical Engineering*. 2023;17:136-152.
2. Chotwanvirat P, Prachansuwan A, Sridonpai P, Kriengsinyos W. Advancements in using AI for dietary assessment based on food images: Scoping review. *Journal of Medical Internet Research*. 2024;26:e51432.
3. Dalakleidi KV, Papadelli M, Kapos I, Papadimitriou K. Applying image-based food-recognition systems on dietary assessment: A systematic review. *Advances in Nutrition*. 2022;13(6):2590-2619.

4. Rao B, Rashid M, Hasan MG, Thunga G. Machine learning in predicting child malnutrition: A meta-analysis of demographic and health surveys data. *International Journal of Environmental Research and Public Health*. 2025;22(3):449.
5. Janssen SM, Bouzembrak Y, Tekinerdogan B. Artificial intelligence in malnutrition: A systematic literature review. *Advances in Nutrition*. 2024;15(9):100264.
6. Parchure P, Besculides M, Zhan S, et al. Malnutrition risk assessment using a machine learning-based screening tool: A multicentre retrospective cohort. *Journal of Human Nutrition and Dietetics* 2024;37(3):622-632.
7. Liao LL, Chang LC, Lai IJ. Assessing the quality of ChatGPT's dietary advice for college students from dietitians' perspectives. *Nutrients*. 2024;16(12):1939.
8. Azimi I, Qi M, Wang L, Rahmani AM, Li Y. Evaluation of LLMs accuracy and consistency in the registered dietitian exam through prompt engineering and knowledge retrieval. *Scientific Reports*. 2025;15(1):1506.
9. Collins GS, Moons KG, Dhiman P, et al. TRIPOD+ AI statement: Updated guidance for reporting clinical prediction models that use regression or machine learning methods. *BMJ*, 2024;385.
10. Liu X, Rivera SC, Faes L, et al. CONSORT-AI and SPIRIT-AI: New reporting guidelines for clinical trials and trial protocols for artificial intelligence interventions. *Investigative Ophthalmology & Visual Science*. 2020;61(7):1617-1617.
11. Rivera SC, Liu X, Chan AW, et al. Guidelines for clinical trial protocols for interventions involving artificial intelligence: The SPIRIT-AI extension. *The Lancet Digital Health*. 2022;2(10):e549-e560.
12. Vasey B, Nagendran M, Campbell B, et al. Reporting guideline for the early stage clinical evaluation of decision support systems driven by artificial intelligence: DECIDE-AI. *BMJ*. 2022;377.
13. Tricco AC, Lillie E, Zarin W, et al. PRISMA extension for scoping reviews (PRISMA-ScR): Checklist and explanation. *Annals of Internal Medicine*, 2018;169(7):467-473.
14. Zeevi D, Korem T, Zmora N, et al. Personalized nutrition by prediction of glycemic responses. *Cell*. 2015;163(5):1079-1094.
15. Berry SE, Valdes AM, Drew DA, et al. Human postprandial responses to food and potential for precision nutrition. *Nature Medicine*. 2020;26(6):964-973.

16. Navarro CLA, Damen JA, Takada T, et al. Systematic review finds “spin” practices and poor reporting standards in studies on machine learning-based prediction models. *Journal of Clinical Epidemiology*. 2023;158:99-110.
17. Nagendran M, Chen, Y, Lovejoy CA, et al. Artificial intelligence versus clinicians: Systematic review of design, reporting standards, and claims of deep learning studies. *BMJ*. 2020;368.
18. Lasko TA, Strobl EV, Stead WW. Why do probabilistic clinical models fail to transport between sites. *NPJ Digital Medicine*. 2024;7(1):53.
19. Wang Y, Song Y, Wang Y, YU L, Wang J. Ethics and governance of artificial intelligence for health: Guidance on large multi-modal models. *Chinese Medical Ethics*. 2024;1001-1022.
20. Obermeyer Z, Powers B, Vogeli C, Mullainathan, S. Dissecting racial bias in an algorithm used to manage the health of populations. *Science*. 2019;366(6464):447-453.
21. Wilkinson MD, Dumontier M, Aalbersberg IJ, et al. The FAIR guiding principles for scientific data management and stewardship. *Scientific Data*. 2016;3(1):1-9.
22. Amugongo LM, Kriebitz A, Boch A, Lütge C. Mobile computer vision-based applications for food recognition and volume and calorific estimation: A systematic review. *In Healthcare*. 2023;11(1):59.
23. Bermingham KM, Linenberg, I, Polidori L, et al. Effects of a personalized nutrition program on cardiometabolic health: A randomized controlled trial. *Nature Medicine*. 2024;30(7):1888-1897.