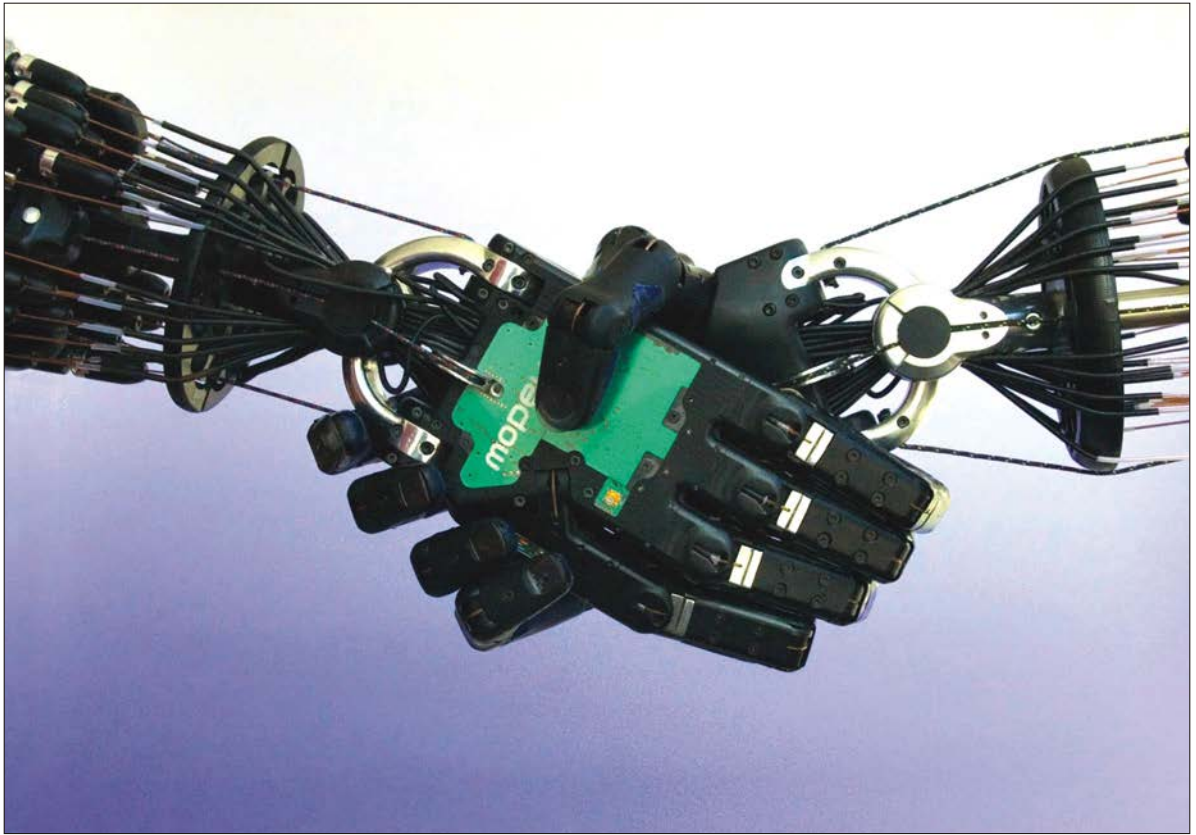




Journal of Artificial Intelligence

ISSN 1994-5450

science
alert

ANSI*net*
an open access publisher
<http://ansinet.com>

Systematic Review

Machine Learning and Metabolomics for Predicting Nutrient and Phytochemical Bioavailability: A Systematic Review

David Chinonso Anih and Kayode Adebisi Arowora

Department of Biochemistry, Faculty of Biosciences, Federal University Wukari, Taraba, Nigeria

Abstract

This systematic review examined the integration of food composition data, processing conditions, *in vitro* digestion, untargeted metabolomics and artificial intelligence for predicting the bioaccessibility and bioavailability of nutrients and phytochemicals in food matrices. The literature search period spanned 2020 to 2025. Studies were retrieved from major biomedical, food science, chemical and computational databases and were eligible when they combined at least two of the three core elements of the review: *in vitro* digestion, untargeted metabolomics and machine learning or artificial intelligence. Following duplicate removal and two-stage screening, 40 studies met the inclusion criteria, with 21 forming the main evidence base for synthesis. The evidence showed that INFOGEST-based static digestion models remain the most widely used framework, although semi-dynamic and dynamic systems offer greater physiological realism. Untargeted liquid chromatography, high-resolution mass spectrometry and gas chromatography mass spectrometry were the principal analytical platforms, but variability in preprocessing, annotation confidence, quality control and reporting standards limited direct comparison across studies. Food processing methods such as heating, fermentation, milling and non-thermal technologies consistently influenced the bioaccessibility of carotenoids, polyphenols, glucosinolates and other bioactive compounds in a matrix-dependent manner. Machine learning approaches, including random forest, gradient boosting, neural networks, graph neural networks and physics-informed models, demonstrated promising capacity to predict bioaccessibility, metabolite transformation and semiquantitative concentration patterns. However, performance frequently declined during external validation, indicating overfitting, instrument effects and domain shift. Overall, the review indicates that artificial intelligence-supported prediction of nutrient and phytochemical bioavailability is scientifically feasible, but the field is not yet fully mature for broad translational use. Progress will depend on standardized benchmark datasets, harmonized digestion and metabolomics protocols, stronger quality control, transparent reporting, independent external validation and wider adoption of FAIR and reproducible workflows. These advances should support more reliable prediction systems for nutrition research, food processing optimization and industrial formulation.

Key words: Bioaccessibility, bioavailability, untargeted metabolomics, *in vitro* digestion, artificial intelligence, infogest

Citation: Anih D.C., and K.A. Arowora, 2026. Machine learning and metabolomics for predicting nutrient and phytochemical bioavailability: A systematic review. J. Artif. Intel., 16: 24-36.

Corresponding Author: David Chinonso Anih, Department of Biochemistry, Faculty of Biosciences, Federal University Wukari, Taraba, Nigeria

Copyright: © 2026 David Chinonso Anih and Kayode Adebisi Arowora. This is an open access article distributed under the terms of the creative commons attribution License, which permits unrestricted use, distribution and reproduction in any medium, provided the original author and source are credited.

Competing Interest: The authors have declared that no competing interest exists.

Data Availability: All relevant data are within the paper and its supporting information files.

INTRODUCTION

Bioaccessibility, bioavailability and bioefficacy are related but distinct concepts that shape how we assess what a food delivers to a living system. Bioaccessibility denotes the fraction of a compound released from the food matrix into the gastrointestinal milieu and therefore available for absorption; bioavailability is the fraction that is actually absorbed and later appears in systemic circulation or target tissues; bioefficacy refers to the physiological effect produced by the absorbed fraction. Because different experimental systems measure different surrogate endpoints (for example, *in vitro* solubilized fraction versus plasma area under the curve), clear operational definitions are critical for comparing studies and for training predictive models that translate *in vitro* signals to likely *in vivo* outcomes^{1,2}.

Phytochemicals and nutrients that commonly appear in predictive workflows include polyphenols (flavonoids, phenolic acids), carotenoids, alkaloids, glucosinolates and small polar metabolites, together with macronutrients such as lipids, proteins and carbohydrates that influence micelle formation, enzymatic accessibility and subsequent transport. Each chemical class brings characteristic physicochemical properties lipophilicity, molecular weight, polarity and common conjugation patterns that shape release during processing, enzymatic transformation in digestion and interactions with bile salts or microbiota. These properties are readily captured as feature patterns by untargeted liquid chromatography-high resolution mass spectrometry (LC-HRMS) and can be paired with metadata on processing and food composition to improve model input fidelity².

Three converging developments motivate integrating *in vitro* digestion, untargeted metabolomics and AI. First, standardized *in vitro* digestion protocols provide repeatable simulations of oral–gastric–intestinal transit and yield bioaccessibility metrics that are amenable to high-throughput experiments, albeit with acknowledged limitations that require careful calibration against *in vivo* data³. Second, untargeted metabolomics produces high-dimensional molecular fingerprints that capture parent compounds and their processing or digestion-derived transformation products; such time-resolved feature matrices are particularly useful as model inputs because they preserve both abundance trends and emergent features not anticipated in targeted assays³. Third, contemporary machine learning approaches, including ensemble learners, gradient boosting and interpretable neural networks, have matured to a point where they can extract

predictive signals from complex metabolomics matrices and rank features that drive outcomes of interest, enabling semi-quantitative predictions of chemical concentrations and transformation pathways during processing and digestion⁴.

Taken together, these elements enable a shift from descriptive characterization (what compounds are present) toward predictive modeling (which compounds will be bioaccessible, absorbed and likely to produce an effect). This review therefore asks: For food matrices or model systems, can AI/ML models trained on combinations of food composition, processing parameters and *in vitro* metabolomics reliably predict bioaccessibility or bioavailability and how well do such predictions concord with *in vivo* or clinical measurements? The review will prioritize studies that combine at least two pillars (*in vitro* digestion, untargeted metabolomics, ML/AI) or use *in vitro*/metabolomics inputs to predict *in vivo* bioavailability and will extract model inputs, architectures, validation strategies and agreement with experimental bioavailability measures to evaluate reproducibility and practical utility⁵.

MATERIALS AND METHODS

Information sources and search strategy: A systematic, multidisciplinary literature search was conducted to identify experimental, analytical and computational studies relevant to food processing, *in vitro* digestion, untargeted metabolomics and machine learning models used to predict bioaccessibility and related endpoints⁶. The search covered major biomedical, chemical, food science, agricultural and machine learning databases, including PubMed/MEDLINE, Scopus, Web of Science, Embase, AGRICOLA, IEEE Xplore, NeurIPS, ICML, arXiv and bioRxiv, together with relevant grey literature such as technical reports and standard operating procedures that support metabolomics quality assurance, quality control and digestion method harmonization^{7,8}. The literature search period spanned 2020 to 2025 and the search terms were adapted to each database and refined iteratively to improve sensitivity while reducing irrelevant results⁸.

Search strategies combined domain terms (bioaccessibility, INFOGEST, *in vitro* digestion) with analytical and modelling terms (untargeted metabolomics, LC-HRMS, LC-QTOF, GC-MS, machine learning, random forest, graph neural network) and were adapted for each database using controlled vocabularies where available⁶. A representative search block used across databases was: (“bioavailability” OR “bioaccessibility”) AND (“*in vitro* digestion” OR INFOGEST)

AND ("untargeted metabolomics" OR "LC-HRMS" OR "LC-QTOF" OR "GC-MS") AND ("machine learning" OR "artificial intelligence" OR "deep learning" OR "random forest" OR "graph neural network"), with iterative refinement to maximize sensitivity while reducing irrelevant retrieval⁷. All search strings, search dates and query logs were archived in a reproducible supplement to this review to support transparency and allow replication of the retrieval process⁸.

Literature search and selection criteria: To identify relevant research at the intersection of food science, analytical chemistry and artificial intelligence, we conducted a systematic, multidisciplinary literature search. Our goal was to capture the full spectrum of work involving food processing, *in vitro* digestion, untargeted metabolomics and the machine learning models used to predict bioaccessibility. We searched major databases, including PubMed/MEDLINE, Scopus, Web of Science, Embase and AGRICOLA. To ensure coverage of the computational side of the field, we also searched IEEE Xplore and major machine learning conference proceedings such as NeurIPS and ICML, while screening preprint servers like arXiv and bioRxiv for recent methodological developments⁹⁻¹¹.

Search strategy: A combination of controlled vocabulary and free-text terms, linked with Boolean operators, was used to connect biological and computational concepts. The search focused on core term combinations such as: "bioavailability" OR "bioaccessibility" AND "*in vitro* digestion" OR "INFOGEST" AND "untargeted metabolomics" OR "LC-HRMS" AND "machine learning" OR "artificial intelligence".

These strings were refined iteratively for each database to ensure high-fidelity results while filtering out irrelevant records. All search logs and dates were archived to maintain a fully reproducible process⁹⁻¹¹.

Inclusion criteria: Studies were eligible for the review if they met the following requirements.

- **Study design:** Empirical research using *in vitro* digestion, untargeted metabolomics, or the development and validation of predictive models
- **Methodological rigor:** Studies had to combine at least two of the three main pillars (digestion, metabolomics, or AI), or use *in vitro*/metabolomics data to predict *in vivo* outcomes
- **Sample type:** Food matrices or model food systems
Outcomes: Clear metrics for bioaccessibility, bioavailability, or model performance, such as R², RMSE, or AUC

- **Reporting threshold:** A clear description of the digestion parameters and the model architecture used⁹⁻¹¹

Exclusion criteria: We excluded records based on the following conditions:

- **Redundancy:** Duplicate records or multiple reports using the same underlying dataset, which were linked to a single primary study
- **Irrelevant scope:** Studies dealing exclusively with mechanical or thermal processing without any digestion or metabolomics component

Insufficient data: Articles lacking specific digestion parameters (enzyme units, pH, time) or those that did not provide enough detail on model validation.

Format: Editorials, commentaries, or papers where the full text was unavailable or not in English.

Study selection and results: The selection followed a rigorous two-tier screening process. First, two reviewers independently screened titles and abstracts to remove clearly irrelevant studies. This was followed by a full-text review to apply the eligibility criteria. Any disagreements between reviewers were resolved through consensus or by a third-party adjudicator.

Our initial search across all platforms retrieved 985 records. After removing 142 duplicates, we screened 843 unique titles and abstracts. From this pool, 718 records were excluded for not meeting the core criteria, such as lacking an AI component or using only targeted assays. The remaining 125 full-text articles were assessed for eligibility.

Ultimately, 40 studies were included in the review. Of these, 21 studies formed the core evidence base for the primary synthesis, focusing on hybrid experimental-model outcomes and predictive accuracy. An additional 19 studies were retained as background or methodological references to provide context on standardized digestion protocols such as INFOGEST, QA/QC in metabolomics and foundational machine learning architectures. This selection process is summarized in the PRISMA flow diagram in Fig. 1.

Data extraction-standardized extraction form (variables to capture): We developed and piloted a standardized extraction form to ensure consistent capture of experimental, analytical and modelling metadata across studies and the form was refined iteratively after comparison of pilot extractions¹². Key variables captured included study identification (author, year,

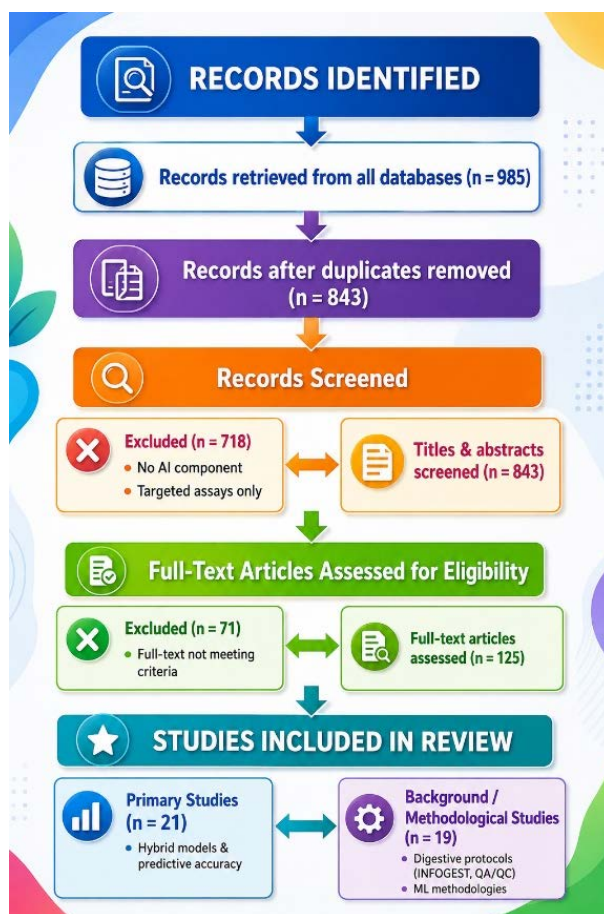


Fig. 1: PRISMA based flow diagram depicting the systematic screening and selection of literature, covering studies published from 2020 to 2025⁹⁻¹¹

A PRISMA flow diagram showing the selection of studies from identification through inclusion. Abbreviations: PRISMA, Preferred Reporting Items for Systematic Reviews and Meta-Analyses; n, number of records/studies

DOI), food matrix and processing details (food type, processing method and parameters), *in vitro* digestion model parameters (INFOGEST static or dynamic, enzyme activities, pH profiles, transit times and sample to fluid ratios), metabolomics platform and acquisition settings (LC-HRMS, GC-MS, ionization mode, acquisition strategy) and feature annotation approach with confidence levels and libraries used¹³.

The extraction form also recorded data processing and QC practices (peak picking, alignment, normalization, pooled QC frequency, internal standards and acceptable RSD thresholds), model inputs and architecture (raw metabolomic features, engineered variables, processing metadata, model family, hyperparameters and feature selection methods), model training strategy (cross-validation scheme, external validation) and performance metrics reported (RMSE, MAE, R^2 , AUC) as well as sample size and replication structure¹³. Reproducibility indicators were explicitly captured, including data and code

availability, adherence to reporting checklists such as TRIPOD or TRIPOD-AI where applicable and whether raw spectral data had been deposited in public repositories¹². Extraction was performed by one reviewer and independently checked by a second reviewer, with disagreements resolved by discussion and, where necessary, by contacting study authors for missing details¹³.

Risk of bias and reporting quality assessment: Because the included literature spans experimental digestion studies and computational prediction work, we used a dual assessment approach: Empirical studies were assessed using an adapted STROBE/QUADAS framework emphasizing pre-analytical factors and analytical QA/QC, while prediction-model studies were assessed using TRIPOD and PROBAST principles and the emerging PROBAST-AI guidance where relevant¹⁴. For empirical digestion and metabolomics studies the emphasis was on sample handling and storage, explicit reporting of

digestion parameters (enzyme units, pH profiles, transit times), analytical validation (internal standards, pooled QC strategy), feature annotation confidence levels and the presence of biological and technical replication to evaluate robustness and precision of measured outcomes¹⁵.

Prediction-model assessment focused on reporting completeness (clear model specification, hyperparameter lists, training/validation splits), risk of bias in predictor measurement and outcome definition, handling of missing data, strategies to reduce overfitting (regularization, cross-validation) and the presence of independent external validation to assess generalizability¹⁶. We produced a reporting quality matrix that maps TRIPOD items against each model study and a risk-of-bias summary using PROBAST domains; studies judged to be at high risk of bias in major domains were flagged for sensitivity analyses and down-weighting in synthesis to protect against misleading pooled estimates¹⁴.

Data synthesis and meta-analysis plan: Our synthesis strategy combined qualitative structured summaries with quantitative meta-analysis when comparable effect measures permitted pooling¹⁷. Qualitative synthesis used structured summary tables (Table 1-6) that present study ID, food matrix, digestion parameters, metabolomics platform and QC indicators, ML model family and performance metrics and data/code availability to facilitate cross-study comparison and subgrouping by phytochemical class, processing type and model family¹⁸. Where studies reported comparable effect measures, such as percent change in bioaccessibility for a target compound after processing or continuous predictive error metrics for the same endpoint, we planned meta-analysis using random-effects models to account for heterogeneity across study designs and measurement platforms¹⁸.

For experimental bioaccessibility outcomes reported as concentrations or percent release we planned to convert reported metrics to standardized mean differences or other common effect sizes with inverse-variance weighting for pooling when appropriate; for predictive model performance, comparable continuous error metrics (RMSE, MAE) were considered for pooling after standardization to a common unit or via normalized error transforms when necessary¹⁷. Where variance estimates were missing, we followed recommended imputation or approximation methods and recorded sensitivity of pooled estimates to these assumptions and where heterogeneity permitted, we used meta-regression to explore methodological moderators such as digestion

protocol (static vs dynamic), metabolomics platform and model family¹⁸. All quantitative analyses were implemented using established statistical packages and custom scripts and transformation equations and assumptions were documented to support reproducibility¹⁸.

Sensitivity analyses, subgroup analyses and publication-bias assessment: We prespecified sensitivity analyses to test robustness of inferences and to probe the impact of methodological choices and study quality on pooled results¹⁹. Primary sensitivity checks included exclusion of studies judged to be at high risk of bias, restricting synthesis to studies using dynamic digestion models and limiting analyses to predictive models that had independent external validation sets to test generalizability¹⁹.

Subgroup analyses were planned by biologically and methodologically meaningful strata including phytochemical class (polyphenols, carotenoids, alkaloids), food matrix (plant, dairy, meat), processing type (thermal, fermentation, milling), metabolomics platform (LC-HRMS vs GC-MS) and model family (tree-based ensembles, deep learning, graph neural networks, physics-informed models) to identify patterns of effect modification and to highlight areas where methodological standardization could improve comparability¹⁹. Publication bias and small-study effects were assessed for any meta-analysed outcome with at least ten studies using funnel plots and Egger's test and trim-and-fill or other sensitivity approaches were used where asymmetry suggested potential reporting bias¹⁹.

RESULTS AND DISCUSSION

Evidence synthesis: Study characteristics and quality: We screened recent work combining *in vitro* digestion plus metabolomics, model development (statistical or mechanistic) and hybrid experimental-model studies to derive a pragmatic typology and quality assessment. The evidence base is dominated by three study types: (i) Empirical *in vitro* digestion coupled with untargeted or targeted metabolomics, typically using INFOGEST derived static protocols or modified dynamic simulators; (ii) Purely computational/modeling research that constructs predictive models of bioaccessibility or MS response from chemical descriptors and instrument metadata; and (iii) Hybrid experimental-model investigations that generate *in vitro* metabolomic profiles and then train machine learning models (e.g., RF/XGBoost/NN) to predict transformation outcomes across processing or digestion conditions. The relative distribution in recent reviews shows

Table 1: Study types, sample sizes and typical quality markers (compact overview)

Study type	Typical sample/replicate size	Key methodological markers	Citation(s)
Empirical (<i>in vitro</i> digestion+metabolomics)	3-6 biological replicates are typical; multiple digestion replicates are common	INFOGEST static/dynamic protocol stated; variable QC reporting (pooled QCs, internal standards)	Reboredo-Rodríguez <i>et al.</i> ²⁰
Model development (<i>in silico</i>)	Datasets range widely (50-10k compounds); often, public chemical libraries	Descriptor sets, cross validation, occasionally external test sets	Reboredo-Rodríguez <i>et al.</i> ²⁰
Hybrid experimental model	Moderate N (20-200 features from experiments)	Explicit feature engineering from metabolites; some use semi quantitation via ML	Reboredo-Rodríguez <i>et al.</i> ²⁰

Abbreviations: INFOGEST, standardized *in vitro* digestion protocol, QC: Quality control, ML: Machine learning and MS: Mass spectrometry

Table 2: Protocol comparison: INFOGEST, static, semi dynamic and dynamic digestion methods

Protocol family	Common enzyme activities and steps	Typical strengths	Typical limitations	Citation(s)
INFOGEST static	Fixed protease/pepsin/trypsin units, oral/gastric/intestinal phases	High inter lab comparability when fully followed	Sensitive to small reagent modifications; micelle formation depends on bile conc.	Tan <i>et al.</i> ²¹
Semi dynamic	Time staggered gastric emptying; variable pH ramps	Better mimicry of gastric emptying; useful for mixed meals	Less standardized; more sample handling	Tan <i>et al.</i> ²¹
Dynamic simulators	Peristaltic pumps, dynamic pH control	Closest to <i>in vivo</i> kinetics	Costly, complex, low throughput.	Kondrashina <i>et al.</i> ²²

Abbreviations: INFOGEST, standardized *in vitro* digestion framework

that empirical digestion+metabolomics remains the largest class, but hybrid experimental-model studies are the fastest growing segment as groups pursue semi quantitative MS outputs and generalizable ML models²⁰.

Quality assessment across papers is heterogeneous. A minority of studies declare full QC procedures (pooled QC, internal standards, injection order checks) and deposit raw data in public repositories; many report only qualitative trends or semi quantitative fold changes without clear limits of detection or reporting on matrix effects²⁰. Hybrid model studies vary in methodological transparency: Some include clear training/test splits and hyperparameter searches, while others rely on single dataset cross validation making external generalizability uncertain²⁰.

Table 1 summarizes the major study types identified in the review, along with their typical sample sizes and key quality indicators.

***In vitro* digestion protocols and standardization issues:**

There are three common families of *in vitro* digestion methods in the literature: INFOGEST derived static methods, semi dynamic adaptations (e.g., timed gastric emptying simulation) and fully dynamic, multi compartmental simulators that better mimic peristalsis, pH gradients and transit times. INFOGEST static approaches provide the greatest cross study comparability because of standardized reagents, enzyme activities and pH/time steps; nevertheless, many groups adapt reagent concentrations (e.g., bile salts), enzyme activities, or incubation times to fit special matrices (lipid rich foods, viscous hydrogels), which produces systematic heterogeneity in reported bioaccessibility²¹.

Semi dynamic methods attempt to capture gastric emptying and bolus dilution effects (useful for lipid digestion

and mixed meals) and reduce the mismatch with *in vivo* transit kinetics, but they require specialized equipment and more complex sample handling; as a result they are less commonly used and less standardized across labs²¹. Static INFOGEST is widely used for polyphenol and phytochemical screening, yet even within INFOGEST variants groups differ on final bile concentrations and whether they apply an oral phase that includes α amylase both materially affecting micelle formation and small molecule solubilization²².

Methodological consequences for metabolite transformation reporting are significant. Shorter gastric incubation times and lower bile concentrations typically under estimate solubilized lipophilic phytochemicals (e.g., carotenoids), while higher protease activity accelerates peptide/protein breakdown and can release bound small molecules that would otherwise remain matrix associated. Transparent reporting of enzyme activities (U/mL), pH, temperature and bile: Sample ratios is therefore essential to interpret or model transformation outcomes.

Table 2 compares the main *in vitro* digestion protocol families and their common strengths and limitations.

Metabolomics platforms, feature processing and annotation challenges:

Two complementary analytical platforms dominate: Liquid Chromatography High Resolution Mass Spectrometry (LC-HRMS) for semi polar and polar non volatile phytochemicals and gas chromatography-mass spectrometry (GC-MS) for volatile and derivatizable small molecules. The LC-HRMS provides the broadest coverage and highest structural hypothesis power (accurate mass, MS/MS), but annotation rates remain modest in untargeted studies (often <30% confidently identified) owing to limited spectral reference libraries and isomeric complexity²³.

Table 3: Platform comparison and annotation considerations in untargeted metabolomics

Platform	Typical preprocessing steps	Annotation rate (typical)	QC/IS practice	Citations(s)
LC-HRMS (UHPLC-QTOF/Orbitrap)	Peak detection, deconvolution, alignment, isotope/adduct grouping	Putative IDs often 20–40% (level 2/3), confirmed IDs <15%	Pooled QC, some labeled IS, retention time standards	Shang <i>et al.</i> ²³
GC-MS (derivatisation)	Deconvolution, library match (NIST), retention index	Higher confident ID for volatiles; >50% for many studies	Internal standards, retention index matching	Shang <i>et al.</i> ²³
Hybrid/strategies	Dual platform fusion (LC+GC), IMS addition	Much broader metabolome coverage	Use of orthogonal standards recommended	Du <i>et al.</i> ²⁴

Abbreviations: LC HRMS, liquid chromatography high resolution mass spectrometry, GC-MS: Gas chromatography mass spectrometry, UHPLC QTOF, ultra high performance liquid chromatography quadrupole time of flight, Orbitrap, high resolution mass analyzer, QC: Quality control and IS: Internal standard

Table 4: Typical quantitative direction and magnitude of processing effects on phytochemical bioaccessibility

Phytochemical class	Typical processing effect (direction)	Representative magnitude change (empirical ranges)	Citation(s)
Carotenoids (lipophilic)	Heat+oil/milling → ↑ bioaccessibility	+20% to +200% (matrix & treatment dependent)	Arfaoui ²⁵
Polyphenols (water soluble)	Boiling → ↓ (leaching), Fermentation → ↑ (deglycosylation)	–30% to +100% depending on method	Li <i>et al.</i> ²⁶
Glucosinolates	Blanching/boiling → ↓ (loss); controlled steaming → ↑ retention	–40% to +10%	Narra <i>et al.</i> ²⁷

Data preprocessing choices (peak picking, deconvolution, alignment, adduct grouping) materially change the feature table: Different software (e.g., MZmine, XCMS, commercial vendor packages) produce non identical feature sets even from the same raw files, making cross study meta analysis challenging. The QC workflows that include pooled QCs, retention time locking with standards, isotopically labeled internal standards for key compound classes and randomized injection orders are increasingly recommended but not universally applied²³.

Semi quantitative reporting is common: Many untargeted pipelines present fold change or peak area ratios rather than absolute concentrations. Recently, groups have coupled machine learning (ionization efficiency models) to enable semi quantitative concentration estimation from LC-HRMS signals this reduces dependence on pure standards but adds model driven uncertainty that must be reported²⁴.

Table 3 outlines the main metabolomics platforms used in the review and the practical issues affecting preprocessing and annotation.

Food processing effects on phytochemical/nutrient bioaccessibility (empirical evidence):

Empirical evidence consistently shows that processing alters bioaccessibility in ways that depend on both compound class and matrix. Thermal treatments produce mixed effects: heat can break cell walls and increase the extractability and micellization of lipophilic compounds such as carotenoids (sometimes increasing bioaccessibility by tens of percent), while water based boiling typically causes leaching losses of water soluble polyphenols and vitamins²⁵.

Fermentation commonly increases bioaccessibility of polyphenols and some carotenoids via enzymatic deglycosylation and microbe mediated breakdown of cell

walls. Numerous *in vitro* digestion+metabolomics studies report increases in accessible aglycone concentrations and higher post digest antioxidant activity after lactic fermentation reported relative bioaccessibility gains often in the range ~10-100% for targeted phenolics depending on strain and matrix²⁶.

Comminution/milling increases the surface area and can substantially increase release of bound phytochemicals; for some matrices, milling plus thermal treatment produces synergistic increases in bioaccessibility (e.g., milled carotenoid rich fruits combined with low temperature heating show higher micellization than either treatment alone). Non thermal technologies (HPP, PEF) are increasingly reported to preserve thermolabile phytochemicals while improving extractability; HPP studies report bioaccessibility increases ~5–35% for phenolics in some juices²⁷.

Table 4 summarizes the typical direction and approximate magnitude of processing effects on major phytochemical classes.

AI/ML model approaches used to predict bioavailability/bioaccessibility:

Machine learning model families used in the field include tree based ensembles (random forest, XGBoost), kernel methods (SVM) and a growing set of neural network architectures ranging from conventional feed forward NNs to graph neural networks (GNNs) for molecular structure representation and physics informed neural networks (PINNs) for coupling data with known physical laws²⁸. Model inputs vary: Molecular descriptors and fingerprints (for individual compound predictions), chromatographic/MS descriptors (for semi quantitative estimation of concentration from signal intensities) and process/digestion parameters (enzyme activities, pH, time) when entire food matrices are modeled.

Table 5: Machine learning model families, inputs and typical use cases

Algorithm family	Typical inputs/features	Common use case	Citation(s)
Random forest/XGBoost	Molecular descriptors, processing parameters	Feature importance, tabular prediction of bioaccessibility	Lutchyn <i>et al.</i> ²⁸
Neural networks/GNNs	Graph representations, learned embeddings	Molecular property prediction, structure-bioavailability	Palm and Krueve ²⁹
PINNs	Same+physical constraints (mass balance)	Kinetic/diffusion coupled digestion modeling	Farea <i>et al.</i> ³¹
MS based IE prediction models	LC/ESI descriptors, eluent, fragment info	Semi quantitation in untargeted LC-HRMS	Beck <i>et al.</i> ³⁰

Abbreviations: RF: Random forest, XGBoost: Extreme gradient boosting, NN: Neural network, GNN: Graph neural network, PINN: Physics informed neural network and IE: Ionization efficiency

Table 6: Performance metrics and validation practices in predictive modeling

Metric/practice	Observed ranges	Common shortcomings	Citation(s)
Internal CV R ² /RMSE	R ² often 0.6-0.95 (task dependent)	Over optimistic without external test	Farea <i>et al.</i> ³¹
External test performance	Often 10-50% worse than CV	Sparse use of external matrices/instruments	Shah <i>et al.</i> ³²
Transfer learning/calibration	Improves cross platform transfer	Requires small labelled calibration sets	Bieber <i>et al.</i> ³³

Abbreviations: CV: Cross validation, RMSE: Root mean square error, MAE: Mean absolute error and R2: Coefficient of determination

Several hybrid pipelines concatenate metabolomics feature vectors with processing metadata and use embedding layers to capture interactions before feeding supervised learners²⁹.

Feature engineering strategies are central. For molecular level prediction the models most often use 1D/2D descriptors (PaDEL, RDKit) or learned graph representations (GNNs) to capture functional group effects relevant to absorption and ionization. For MS semi quantitation, descriptor sets include retention time, MS1 peak shape characteristics, MS/MS fragmentation patterns and instrument/eluent conditions; supervised models are then trained to map those descriptors to ionization efficiency or a relative response factor enabling concentration estimates without compound standards³⁰.

Emerging approaches show promise: (i) PINNs have been proposed as a way to impose conservation laws or mass balance constraints (for digestion kinetics) inside learning systems, improving extrapolation across conditions; (ii) Models that predict ionization efficiency or relative response factors from MS descriptors permit semi quantitative interpretation of untargeted data and are rapidly being incorporated into metabolomics to model pipelines³¹.

Table 5 in above paragraph summarizes the principal artificial intelligence and machine learning approaches used to predict bioaccessibility and related outcomes.

Model performance and external validation: Model performance reporting is inconsistent. Many studies report internal cross validation performance (R², RMSE, AUC) and hyperparameter selection; however, far fewer include independent external test sets that reflect distinct food matrices or instruments. When external validation is performed, performance typically drops compared with internal cross validation sometimes substantially indicating overfitting to matrix specific patterns and an underappreciated role of instrument/matrix covariates³².

For ionization efficiency and semi quantification models, published RMSEs for log IE predictions commonly fall in the ~0.5-1.0 log units range on held out test sets if models are trained on broad, instrument diverse datasets; however, transfer between chromatographic systems is nontrivial and often requires calibration or transfer learning to preserve accuracy³³.

Generalizability across food matrices remains the major challenge. Models trained on aqueous extracts or simple matrices underperform when applied to complex emulsions, fiber rich solids, or samples with heavy protein/lipid content because those matrices alter extraction efficiency and MS response. Best practice examples combine stratified external validation (hold out matrices), measurement of matrix effects and reporting of limits of applicability³³.

Table 6 presents the common validation approaches and performance metrics reported in prediction studies, along with recurring shortcomings.

Mechanistic insights and metabolite transformations revealed by untargeted metabolomics:

Untargeted metabolomics across digestion and fermentation experiments has revealed recurrent transformation patterns that are mechanistically interpretable and that also serve as predictive features for ML models. The most common patterns are: (i) Glycoside hydrolysis/deglycosylation of flavonoids and other glycosylated polyphenols to yield aglycones (which are more readily absorbed in the small intestine); (ii) Oxidative modifications producing quinone like products or phenolic oxidation derivatives; and (iii) Conjugation reactions (sulfation, glucuronidation in *ex vivo* systems or phase I/II metabolite modelling) during simulated intestinal phases or during microbial fermentation³⁴.

Digestion studies repeatedly show loss of polymeric or bound phenolics and a corresponding rise in smaller phenolic acids after gastric/intestine steps or after colonic fermentation

Table 7: Recurrent metabolite transformation motifs and modelling implications

Transformation	Typical detection pattern	Modelling features that capture it	Citation(s)
Deglycosylation (glycoside → aglycone)	Loss of sugar moiety mass; rise in aglycone MS1 and fragment ions	Glycosylation flags, MW delta, predicted glycosidase sites	Bieber <i>et al.</i> ³³
Oxidation	+O mass shifts; formation of quinone fragments	Number of reactive hydroxyls, redox propensity descriptors	Ziólkiewicz <i>et al.</i> ³⁴
Conjugation (sulfate/glucuronide)	+80/+176 mass additions in LC-HRMS	Conjugation susceptibility descriptors, pKa	Yang <i>et al.</i> ³⁵

Abbreviations: MS1: First stage mass spectrometry, MW: Molecular weight, pKa: Acid dissociation constant and pKb: Base dissociation constant

Table 8: Example pipeline steps and primary bottlenecks

Pipeline stage	Common best practice	Typical bottleneck	Citation(s)
Experimental design	Pre registered sample metadata; replicates	Missing enzyme/bile reporting	Yang <i>et al.</i> ³⁵
Acquisition and QC	Pooled QCs, randomized runs, IS	Incomplete IS panels	Bremer and Fiehn ³⁶
Preprocessing	Use of documented software with parameters	Software heterogeneity → different feature sets	Zulfiqar <i>et al.</i> ³⁷
Modeling	External test sets; uncertainty estimates	Small datasets & domain shift	Yang <i>et al.</i> ³⁵ , Bremer and Fiehn ³⁶ and Zulfiqar <i>et al.</i> ³⁷

Abbreviations: FAIR: Findable, accessible, interoperable, reusable, SMetaS: Sample metadata standardizer, QC: Quality control and IS: Internal standard

in fecal inocula. Microbial metabolism and bacterial glycosidases explain many deglycosylation events detected post fermentation and such transformations are strong predictors of increases in solubility and Short Chain Fatty Acids (SCFAs) production in subsequent fecal fermentation assays³⁵.

These mechanistic transformations have been successfully encoded into predictive features: Presence/absence of glycosylation motifs, log P, pKa/pKb estimates and predicted enzyme cleavage sites are among the highest ranked predictors when models learn to forecast whether a parent phytochemical will increase or decrease in bioaccessible fraction after fermentation/digestion³⁵.

Table 7 captures the major metabolite transformation patterns observed during digestion and fermentation and their relevance for model development.

Integration pipelines: How studies combined *in vitro* results, metabolomics and ML (workflows and bottlenecks):

Typical integration pipelines follow this backbone: Controlled *in vitro* digestion/processing experiment→sample clean up and LC-HRMS (±GC-MS) acquisition with pooled QC→raw data preprocessing and feature table generation→feature annotation (putative ID) and internal standard normalization→feature engineering (molecular descriptors, processing metadata, MS descriptors)→supervised ML modeling and validation→biological interpretation and (occasionally) experimental follow up. Pipelines that implement FAIR workflows, use persistent identifiers and produce reproducible notebooks are still the minority, but workbench initiatives and workflow packages are emerging³⁵.

Key bottlenecks are: Inconsistent metadata reporting (enzyme units, bile concentration, sample:enzyme ratio), inconsistent QC and IS usage across labs, heterogeneity of preprocessing pipelines producing incompatible feature

tables and the limited size/variety of public training sets for ML tasks. Tools such as SMetaS (sample metadata standardizer) and FAIR workflow packages have been proposed and piloted to address these barriers³⁶.

Best practices emerging across successful pipelines include: (1) Built in QC at acquisition and preprocessing stages (pooled QCs, blanks and IS); (2) Depositing raw data and metadata in standard repositories; (3) Reporting limits of detection and semi quantitative error estimates for ML derived concentrations; and (4) Including independent external validation matrices in model evaluation³⁷.

Table 8 outlines the major workflow stages in integrated experimental and computational pipelines, together with the most common bottlenecks.

Sources of heterogeneity, limitations in the evidence base and bias patterns:

Heterogeneity stems from many sources: Digestion model variants (static vs dynamic vs semi dynamic), inconsistent reporting of enzyme units and bile concentrations, instrument differences (different ion sources, collision energies), variable use of internal standards and post processing choices (peak picking thresholds, adduct grouping). These technical heterogeneities are compounded by biological heterogeneity (differences in sample matrices, cultivar/variety, processing recipes) and study design (sample sizes, replication strategy)³⁷.

Limitations include the small size of many experimental datasets, infrequent external validation of models, limited availability of public raw data associated with full metadata and a tendency for positive result bias where studies highlight compounds that increase in accessibility while under reporting compounds that decrease or disappear post processing. Together these issues reduce reproducibility and inflate apparent effect sizes in narrative syntheses³⁷.

Table 9: Sources of heterogeneity and recommended mitigation strategies

Source of heterogeneity	Effect on results	Mitigation recommendations	Citation(s)
Digestion protocol variants	Shifts in micellization and solubilization	Report full protocol; adopt community standards	Zulfiqar <i>et al.</i> ³⁷
Instrument/platform differences	Different ionization efficiencies	Multi instrument datasets; transfer learning	Zulfiqar <i>et al.</i> ³⁷
Annotation uncertainty	Uncertain features – model noise	Provide confidence levels; use multiple orthogonal IDs	Hajjar <i>et al.</i> ³⁸

Abbreviations: No table specific abbreviations are used beyond standard methodological terms

Table 10: Prioritized research agenda (short term–long term)

Priority	Action	Expected impact	Citation(s)
1 (short)	Create and publish benchmark digestion+MS datasets w/ metadata	Enables fair model comparison	Niehues <i>et al.</i> ³⁹
2 (short mid)	Develop standard semi quantitative ML pipelines for IE prediction	Broader, cheaper screening	Pang <i>et al.</i> ⁴⁰
3 (mid)	Implement PINNs for digestion kinetics	Improved extrapolation and interpretability	Niehues <i>et al.</i> ³⁹
4 (long)	Multi lab external validation consortium	Industry translation, regulatory confidence	Pang <i>et al.</i> ⁴⁰

Abbreviations: PINN: Physics informed neural network, FAIR: Findable, accessible, interoperable, reusable and MS: Mass spectrometry

Bias patterns to watch for are: (i) Instrument bias models trained on one instrument or lab do not generalize to others; (ii) Matrix bias models trained on aqueous extracts fail on complex solid matrices; and (iii) Annotation bias low annotation rates and tentative IDs lead to model features built on uncertain assignments. Transparent reporting, public data deposition and challenge datasets are required to reveal and correct these biases³⁸.

Table 9 summarizes the major sources of heterogeneity in the evidence base and the strategies recommended to reduce bias and improve comparability.

Research gaps, future directions and translational potential: Based on the synthesis above, priority research directions are:

- Standardized benchmark datasets curated, multi matrix digestion+LC-HRMS+GC-MS datasets with full metadata (enzyme units, bile concentrations, instrument parameters) and standardized QC measures. These would enable robust benchmarking of preprocessing and modeling approaches and reduce cross study heterogeneity³⁹
- Physics informed and hybrid ML methods-PINNs that enforce mass balance and kinetic constraints in digestion models can improve extrapolation across time and conditions; hybrid approaches that combine domain knowledge (e.g., glycosidase activity rules) with data driven learners will enhance interpretability and reduce data requirements³⁹
- Transfer learning and calibration strategies-transfer learning for instrument/matrix adaptation and small calibration experiments to correct for platform differences (i.e., few shot calibration) will be essential for model deployment across labs and food industries⁴⁰

- Synthetic data and augmentation simulated chromatograms and in silico fragmentation (combined with generative approaches) could augment scarce labeled datasets for rare phytochemicals; however, synthetic data must be carefully validated against real measurements to avoid introducing bias⁴⁰
- Community standards and FAIR pipelines metadata standardizers (SMetaS), FAIR workflows and common file formats for semi quantitative results (with uncertainty reporting) are required to move the field from isolated studies to multi center reproducible science⁴⁰

Translational potential is real: Semi quantitative pipelines that estimate concentrations from untargeted data via ML could drastically reduce analytical costs for industrial screening and nutritional assessment and PINN enabled kinetic models could power *in silico* formulations that optimize processing to enhance bioaccessible fractions. To reach that potential the field must invest in benchmark resources, multi site validation challenges and transparent reporting frameworks.

Table 10 presents the recommended research priorities, expected impact and implementation sequence for advancing the field.

CONCLUSION

This review shows that the combined use of *in vitro* digestion, untargeted metabolomics and machine learning offers a strong and timely pathway for predicting nutrient and phytochemical bioaccessibility. The evidence indicates that food composition and processing can markedly influence the release, transformation and eventual availability of bioactive compounds.

Among the digestion models reviewed, INFOGEST based systems remain the most widely adopted, although semi dynamic and dynamic approaches may better reflect physiological conditions. Untargeted metabolomics has expanded the ability to capture complex biochemical changes during digestion and processing, but annotation uncertainty and inconsistent quality control still limit comparability. Machine learning models, especially ensemble methods and neural network based approaches, have shown encouraging predictive capacity, yet their performance often weakens under external validation. This pattern suggests that stronger benchmarking, standardized metadata reporting and broader cross platform validation are essential before routine translation can be achieved. Future studies should prioritize harmonized digestion protocols, richer metabolomics reference datasets and transparent reporting of preprocessing and model development steps. There is also a clear need for physics informed and hybrid models that combine biochemical knowledge with data driven learning to improve interpretability and generalizability. Building multi laboratory datasets and FAIR compliant workflows will be crucial for moving the field from isolated studies to reproducible predictive science. Overall, the review confirms that AI supported food bioavailability prediction is promising, but its full potential will depend on standardization, validation and coordinated future research.

SIGNIFICANCE STATEMENT

This manuscript demonstrates that integrating *in vitro* digestion, untargeted metabolomics and machine learning creates a powerful framework for predicting the bioaccessibility and bioavailability of nutrients and phytochemicals in processed foods. It also shows that while current models are promising, their reliability is constrained by inconsistent digestion protocols, variable metabolomics quality control and limited external validation. The study, therefore, highlights the need for standardized datasets, harmonized workflows and stronger validation strategies to advance reliable and reproducible food bioavailability prediction.

ACKNOWLEDGMENT

The authors sincerely acknowledge the valuable scholarly contributions that informed this systematic review and the efforts of all researchers whose published work formed the foundation of this synthesis. We also appreciate the careful methodological guidance that supported the development of this manuscript and strengthened its overall quality.

REFERENCES

1. Page, M.J., J.E. McKenzie, P.M. Bossuyt, I. Boutron and T.C. Hoffmann *et al.*, 2021. The PRISMA 2020 statement: An updated guideline for reporting systematic reviews. *BMJ*, Vol. 372. 10.1136/bmj.n71.
2. Galal, A., M. Talal and A. Moustafa, 2022. Applications of machine learning in metabolomics: Disease modeling and classification. *Front. Genet.*, Vol. 13. 10.3389/fgene.2022.1017340.
3. Vital, N., A.C. Gramacho, M. Silva, M. Cardoso and P. Alvito *et al.*, 2024. Challenges of the application of *in vitro* digestion for nanomaterials safety assessment. *Foods*, Vol. 13. 10.3390/foods13111690.
4. Zhang, J.D., C. Xue, V.B. Kolachalama and W.A. Donald, 2023. Interpretable machine learning on metabolomics data reveals biomarkers for Parkinson's disease. *ACS Cent. Sci.*, 9: 1035-1045.
5. Zhong, P., X. Wei, X. Li, X. Wei and S. Wu *et al.*, 2022. Untargeted metabolomics by liquid chromatography-mass spectrometry for food authentication: A review. *Compr. Rev. Food Sci. Food Saf.*, 21: 2455-2488.
6. Mulet-Cabero, A.I., L. Egger, R. Portmann, O. Ménard and S. Marze *et al.*, 2020. A standardised semi-dynamic *in vitro* digestion method suitable for food-an international consensus. *Food Funct.*, 11: 1702-1720.
7. Zhou, H., Y. Tan and D.J. McClements, 2023. Applications of the INFOGEST *in vitro* digestion model to foods: A review. *Annu. Rev. Food Sci. Technol.*, 14: 135-156.
8. Xavier, A.A.O. and L.R.B. Mariutti, 2021. Static and semi-dynamic *in vitro* digestion methods: State of the art and recent achievements towards standardization. *Curr. Opin. Food Sci.*, 41: 260-273.
9. Damen, J.A.A., K.G.M. Moons, M. van Smeden and L. Hooft, 2023. How to conduct a systematic review and meta-analysis of prognostic model studies. *Clin. Microbiol. Infect.*, 29: 434-440.
10. Rayner, D.G., B. Kim and F. Foroutan, 2024. A brief step-by-step guide on conducting a systematic review and meta-analysis of prognostic model studies. *J. Clin. Epidemiol.*, Vol. 170. 10.1016/j.jclinepi.2024.111360.
11. Collins, G.S., K.G.M. Moons, P. Dhiman, R.D. Riley and A.L. Beam *et al.*, 2024. TRIPOD+AI statement: Updated guidance for reporting clinical prediction models that use regression or machine learning methods. *BMJ*, Vol. 385. 10.1136/bmj-2023-078378.
12. Moons, K.G.M., J.A.A. Damen, T. Kaul, L. Hooft and C.A. Navarro *et al.*, 2025. PROBAST+AI: An updated quality, risk of bias, and applicability assessment tool for prediction models using regression or artificial intelligence methods. *BMJ*, Vol. 388. 10.1136/bmj-2024-082505.

13. Chi, J., J. Shu, M. Li, R. Mudappathi and Y. Jin *et al.*, 2024. Artificial intelligence in metabolomics: A current review. *TrAC Trends Anal. Chem.*, Vol. 178. 10.1016/j.trac.2024.117852.
14. González-Domínguez, Á., N. Estanyol-Torres, C. Brunius, R. Landberg and R. González-Domínguez, 2024. *QComics*: Recommendations and guidelines for robust, easily implementable and reportable quality control of metabolomics data. *Anal. Chem.*, 96: 1064-1072.
15. Mosley, J.D., W.B. Dunn, J. Kuligowski, M.R. Lewis and M.E. Monge *et al.*, 2024. Metabolomics 2023 workshop report: Moving toward consensus on best QA/QC practices in LC-MS-based untargeted metabolomics. *Metabolomics*, Vol. 20. 10.1007/s11306-024-02135-w.
16. Lippa, K.A., J.J. Aristizabal-Henao, R.D. Beger, J.A. Bowden and C. Broeckling *et al.*, 2022. Reference materials for MS-based untargeted metabolomics and lipidomics: A review by the metabolomics quality assurance and quality control consortium (mQACC). *Metabolomics*, Vol. 18. 10.1007/s11306-021-01848-6.
17. Broeckling, C.D., R.D. Beger, L.L. Cheng, R. Cumeras and D.J. Cuthbertson *et al.*, 2023. Current practices in LC-MS untargeted metabolomics: A scoping review on the use of pooled quality control samples. *Anal. Chem.*, 95: 18645-18654.
18. Zhang, Y., S. Fan, G. Wohlgemuth and O. Fiehn, 2023. Denoising autoencoder normalization for large-scale untargeted metabolomics by gas chromatography-mass spectrometry. *Metabolites*, Vol. 13. 10.3390/metabo13080944.
19. Reiser, P., M. Neubert, A. Eberhard, L. Torresi and C. Zhou *et al.*, 2022. Graph neural networks for materials science and chemistry. *Commun. Mater.*, Vol. 3. 10.1038/s43246-022-00315-6.
20. Reboredo-Rodríguez, P., L. Olmo-García, M. Figueiredo-González, C. González-Barreiro, A. Carrasco-Pancorbo and B. Cancho-Grande, 2021. Application of the INFOGEST standardized method to assess the digestive stability and bioaccessibility of phenolic compounds from Galician extra-virgin olive oil. *J. Agric. Food Chem.*, 69: 11592-11605.
21. Tan, Y., H. Zhou and D.J. McClements, 2022. Application of static *in vitro* digestion models for assessing the bioaccessibility of hydrophobic bioactives: A review. *Trends Food Sci. Technol.*, 122: 314-327.
22. Kondrashina, A., E. Arranz, A. Cilla, M.A. Faria and M. Santos-Hernández *et al.*, 2024. Coupling *in vitro* food digestion with *in vitro* epithelial absorption; recommendations for biocompatibility. *Crit. Rev. Food Sci. Nutr.*, 64: 9618-9636.
23. Shang, W., G. Wei, H. Li, G. Zhao and D. Wang, 2025. Advances in high-resolution mass spectrometry-based metabolomics: Applications in food analysis and biomarker discovery. *J. Agric. Food Chem.*, 73: 3305-3325.
24. Du, X., F. Dastmalchi, H. Ye, T.J. Garrett and M.A. Diller *et al.*, 2023. Evaluating LC-HRMS metabolomics data processing software using FAIR principles for research software. *Metabolomics*, Vol. 19. 10.1007/s11306-023-01974-3.
25. Arfaoui, L., 2021. Dietary plant polyphenols: Effects of food processing on their content and bioavailability. *Molecules*, Vol. 26. 10.3390/molecules26102959.
26. Li, R., Z. Wang, K.W. Kong, P. Xiang, X. He and X. Zhang, 2022. Probiotic fermentation improves the bioactivities and bioaccessibility of polyphenols in *Dendrobium officinale* under *in vitro* simulated gastrointestinal digestion and fecal fermentation. *Front. Nutr.*, Vol. 9. 10.3389/fnut.2022.1005912.
27. Narra, F., E. Piragine, G. Benedetti, C. Ceccanti and M. Florio *et al.*, 2024. Impact of thermal processing on polyphenols, carotenoids, glucosinolates, and ascorbic acid in fruit and vegetables and their cardiovascular benefits. *Compr. Rev. Food Sci. Food Saf.*, Vol. 23. 10.1111/1541-4337.13426.
28. Lutchyn, T., M. Mardal and B. Ricaud, 2025. Efficient learning of molecular properties using graph neural networks enhanced with chemistry knowledge. *ACS Omega*, 10: 54421-54429.
29. Palm, E. and A. Krueve, 2022. Machine learning for absolute quantification of unidentified compounds in non-targeted LC/HRMS. *Molecules*, Vol. 27. 10.3390/molecules27031013.
30. Beck, A.G., M. Muhoberac, C.E. Randolph, C.H. Beveridge, P.R. Wijewardhane, H.I. Kenttämä and G. Chopra, 2024. Recent developments in machine learning for mass spectrometry. *ACS Meas. Sci. Au*, 4: 233-246.
31. Farea, A., O. Yli-Harja and F. Emmert-Streib, 2024. Understanding physics-informed neural networks: Techniques, applications, trends, and challenges. *AI*, 5: 1534-1557.
32. Shah, A., F. Bi and J. Yang, 2026. Machine learning to predict food effects during drug development: A comprehensive review. *J. Cheminf.*, Vol. 18. 10.1186/s13321-025-01131-z.
33. Bieber, S., T. Letzel and A. Krueve, 2023. Electrospray ionization efficiency predictions and analytical standard free quantification for SFC/ESI/HRMS. *J. Am. Soc. Mass Spectrom.*, 34: 1511-1518.
34. Ziółkiewicz, A., K. Kasprzak-Drozd, A. Wójtowicz, T. Oniszcuk and M. Gancarz *et al.*, 2023. The effect of *in vitro* digestion on polyphenolic compounds and antioxidant properties of sorghum (*Sorghum bicolor* (L.) Moench) and sorghum-enriched pasta. *Molecules*, Vol. 28. 10.3390/molecules28041706.
35. Yang, F., C. Chen, D. Ni, Y. Yang and J. Tian *et al.*, 2023. Effects of fermentation on bioactivity and the composition of polyphenols contained in polyphenol-rich foods: A review. *Foods*, Vol. 12. 10.3390/foods12173315.
36. Bremer, P.L. and O. Fiehn, 2023. SMetaS: A sample metadata standardizer for metabolomics. *Metabolites*, Vol. 13. 10.3390/metabo13080941.

37. Zulfiqar, M., M.R. Crusoe, B. König-Ries, C. Steinbeck, K. Peters and L. Gadelha, 2024. Implementation of FAIR practices in computational metabolomics workflows-a case study. *Metabolites*, Vol. 14. 10.3390/metabo14020118.
38. Hajjar, G., M.C.B. Santos, J. Bertrand-Michel, C. Canlet and F. Castelli *et al.*, 2023. Scaling-up metabolomics: Current state and perspectives. *TrAC Trends Anal. Chem.*, Vol. 167. 10.1016/j.trac.2023.117225.
39. Niehues, A., C. de Visser, F.A. Hagenbeek, P. Kulkarni and R. Pool *et al.*, 2024. A multi-omics data analysis workflow packaged as a FAIR digital object. *GigaScience*, Vol. 13. 10.1093/gigascience/giad115.
40. Pang, Z., L. Xu, C. Viau, Y. Lu, R. Salavati, N. Basu and J. Xia, 2024. MetaboAnalystR 4.0: A unified LC-MS workflow for global metabolomics. *Nat. Commun.*, Vol. 15. 10.1038/s41467-024-48009-6.