



Journal of Advanced Pharmaceutical Sciences and Natural Products (JAPSNP)

ARTIFICIAL INTELLIGENCE FOR PHARMACEUTICAL FORMULATION: OPPORTUNITIES, CHALLENGES, AND BARRIERS TO INDUSTRIAL ADOPTION

Sakshi Shyam Nagrikar*

*Jagdambha Institute of Pharmacy and Research,
Kalamb*

Corresponding author:

Sakshi Shyam Nagrikar

Email: sakshishnagrikar2001@gmail.com

ABSTRACT

Artificial intelligence has moved from a speculative promise to a working method across pharmaceutical formulation, yet its presence in routine industrial development remains thin relative to the volume of published proof of concept. This review reads that gap directly. It first maps where machine learning now delivers genuine value in formulation science: predicting solubility, dissolution, disintegration and physical stability from composition; navigating the excipient design space faster than sequential experimentation allows; generating candidate formulations by inverse design; and closing the loop between prediction and experiment through autonomous, self-driving laboratories. Representative platforms already predict cyclodextrin complexation

and solid dispersion stability with test-set accuracies that would have seemed implausible a decade ago. The review then argues that the constraint on adoption is no longer principally algorithmic. Three coupled barriers dominate. The data problem is structural, because formulation datasets are small, proprietary, fragmented and inconsistently reported, which limits model generalisation more than any shortfall in model architecture. The regulatory position is genuinely evolving rather than settled, and the frameworks now emerging from the United States Food and Drug Administration and the European Medicines Agency ask questions about model lifecycle, credibility and human oversight that most published models were never built to answer. The organisational barrier is the least discussed and possibly the most binding. It reaches across validation under existing quality systems, the scarcity of staff fluent in both formulation and data science, and the difficulty of scaling a pilot into a validated production workflow. The paper concludes that the decisive work now sits downstream of the model, in data infrastructure, benchmarking, regulatory engagement and workforce capability.

Keywords

artificial intelligence; machine learning; pharmaceutical formulation; computational pharmaceutics; drug delivery; quality by design; regulatory science; technology adoption



1. Introduction

Pharmaceutical formulation is the discipline that turns a promising molecule into a medicine a patient can take. It decides whether a poorly soluble compound reaches the systemic circulation, whether a tablet holds together on a high-speed press and then breaks apart in the gut, and whether a product survives two years on a shelf without losing potency. For most of its history the discipline has advanced through structured experimentation, and it moved over time from one-factor-at-a-time screening to statistical design of experiments, with the formulator's accumulated intuition supplying the hypotheses that the experiments then tested. This approach works, but it is expensive in time and material, and it scales poorly against the combinatorial size of the formulation space, where a handful of excipients, ratios and process settings generate more candidate products than any laboratory can prepare.

The economic pressure to do better is not new. The cost of bringing a new medicine to market has risen for decades while the number of approvals per dollar of research spending has fallen, a pattern that has resisted the obvious remedies and that pushed the industry to look for efficiency wherever development effort concentrates.¹ Formulation and process development is one such concentration. It sits on the critical path between a validated candidate and a filed product, and delays there translate directly into lost patent-protected market time.

Against this backdrop, artificial intelligence has been presented as a general accelerant for drug

research. Machine learning has demonstrably reshaped parts of discovery, from target identification to property prediction, and reviews of its application across the pipeline have multiplied.^{2,3} Formulation was slower to attract the same attention, partly because its datasets are smaller and messier than the large chemical libraries that fuelled early discovery models, and partly because its outputs must satisfy a manufacturing and regulatory apparatus that discovery models rarely touch. That situation has changed. A substantial body of work now applies machine learning to preformulation, dosage form design, drug delivery and manufacturing control, and dedicated reviews of the field have begun to appear.^{4,5}

The difficulty is that enthusiasm has outrun deployment. The literature is rich in models that predict a formulation outcome accurately on a held-out test set, and comparatively poor in accounts of such models actually governing a formulation decision inside a regulated development programme. This review takes that discrepancy as its central question. It does not attempt an exhaustive catalogue of every algorithm applied to every dosage form, because several recent surveys already serve that purpose.^{4,5} Instead it argues a position: the technical case for artificial intelligence in formulation is now largely made at the level of demonstration, and the factors holding back industrial adoption are predominantly not the models themselves but the data infrastructure they depend on, the regulatory expectations they must meet, and the organisational systems into



which they must fit. The paper first surveys the opportunities where the method has proven its worth, then examines the technical challenges that limit model performance, and then treats at greater length the regulatory, quality and organisational barriers that determine whether a working model ever becomes a working practice.

2. Scope and approach

This is a narrative review. The literature was identified through structured searches of PubMed, Europe PMC and Google Scholar combining the concepts of artificial intelligence, machine learning, deep learning and computational pharmaceutics with formulation, dosage form, drug delivery, manufacturing and regulatory terms, supplemented by the guidance documents of the principal regulatory agencies and by forward and backward citation tracing from key sources. Emphasis fell on primary studies that report quantitative model performance, on platforms that expose their methods and data, and on official regulatory texts rather than secondary commentary on them. The boundary of the review is the formulation and manufacturing segment of the pipeline, so molecular discovery, clinical trial analytics and pharmacovigilance are cited only where they illuminate the formulation case. The review is deliberately weighted toward the adoption question, which means that a modest technical result with clear implications for deployment is given more room than a stronger result with none.

3. The opportunity landscape

3.1 From experimental design to data-driven formulation

The first and most mature contribution of machine learning to formulation is prediction of a product property from its composition and process, in place of, or ahead of, the experiment that would otherwise measure it. Solubility, the property that governs the fate of the large fraction of modern candidates that are poorly water soluble, was an early target and remains a benchmark. Models trained on solubility across many solvents now predict the outcome with cross-validated coefficients of determination above 0.9, which is accurate enough to triage solvent and cosolvent choices before any material is consumed.⁶ The same logic extends across preformulation, where property prediction narrows a wide initial space to a short list worth making.

What distinguishes the current generation from the expert systems and early neural networks of the 1990s is less the ambition than the breadth and the honesty of validation. Contemporary studies report performance on held-out data, examine which input features drive the prediction, and increasingly release the data and code, which allows the results to be checked and reused. The shift is from a model as a one-off artefact to a model as a reusable instrument, and it is that shift, more than any single accuracy figure, that makes the approach industrially interesting.



3.2 Predicting product performance

The properties that matter most to a formulator are the ones that decide whether a product works and lasts, and machine learning has been applied to each of them. Dissolution and disintegration, the in vitro behaviours that stand proxy for how a solid dosage form releases its drug, have proven tractable. A recent study of tablet disintegration assembled roughly two thousand data points spanning molecular, physical and compositional descriptors and compared several regression approaches, with a sparse Bayesian model reaching a test-set coefficient of determination of 0.96 and a mean absolute percentage error near twelve percent.^{7,8} Beyond the headline accuracy, the analysis identified wetting time as the dominant predictor of disintegration, a result that is mechanistically sensible and that shows how a well-built model can return interpretable formulation knowledge rather than an opaque number.

Physical stability, the tendency of a metastable formulation such as an amorphous solid dispersion to revert over time, is harder because it unfolds over months and resists rapid measurement. Machine learning offers a route around the wait by predicting the eventual outcome from measurable early properties. Models trained on curated stability datasets have classified solid dispersions as stable or unstable with usable accuracy, which turns a slow physical test into a fast computational screen that can rank candidate carriers before the stability chambers are loaded.⁷ The lineage of this idea runs back to early demonstrations that a deep

neural network could predict the in vitro behaviour of cyclodextrin and solid dispersion formulations directly from paired descriptors of drug and excipient, which established that formulation outcomes are learnable from composition at all.⁹ Encapsulation efficiency, particle size and release from more complex carriers have been addressed in the same spirit, which extends property prediction from simple tablets to nanoparticles and liposomes.

The clearest demonstration that these capabilities can be integrated rather than scattered is the emergence of unified prediction platforms. One freely available web platform brings several delivery systems under a single interface, covering cyclodextrin complexes, solid dispersions, phospholipid complexes, nanocrystals, self-emulsifying systems and liposomes, each backed by its own curated dataset and model.⁶ The individual datasets are revealing in scale, ranging from a few hundred records for phospholipid complexes to several thousand for self-emulsifying systems, and the modelling choices are equally revealing: tree-based ensembles and support vector machines were selected over deep neural networks because, with datasets of this size and feature count, the simpler methods generalised better.⁶ The performance such a platform achieves is respectable rather than perfect, and reading it honestly is instructive. Solubility regression reached a test-set coefficient of determination near 0.91, and prediction of the free energy of cyclodextrin complexation reached about 0.87, while the classification tasks distinguishing, for



example, stable from unstable systems exceeded 0.8 accuracy with areas under the curve above 0.9.⁶ Predictions for endpoints backed by the smallest datasets, such as certain nanocrystal size outputs, were visibly weaker, which is the small-data effect made concrete. That single pattern, strong where data are plentiful and weak where they are scarce, encapsulates a truth about the field to which the review returns below.

3.3 Generative and inverse design

Prediction answers a forward question, namely what a given formulation will do. The more ambitious inverse question asks what formulation would achieve a desired target, and this is where generative methods enter. Rather than scoring candidates one by one, a generative model learns the distribution of viable formulations and proposes new ones conditioned on a specification, in principle collapsing the search from enumeration to direct design. The approach is well established in molecular discovery and is migrating toward formulation, where the design variables are excipients, ratios and process settings rather than atoms and bonds. The promise is substantial, because inverse design targets the exact bottleneck that forward screening only mitigates, but the formulation instance is less mature than its discovery counterpart, and the generated candidates still require experimental confirmation. A generative model can propose a formulation that is statistically plausible yet physically unmakeable or unstable, so the output is a hypothesis to be tested rather than an answer to be trusted, and its value depends on how cheaply the tests can be

run. Inverse design is therefore best read at present as a rapidly developing opportunity rather than a settled capability, and one whose practical payoff is tightly coupled to the closed-loop experimentation discussed next.

3.4 Closed-loop and autonomous formulation

The logical endpoint of pairing prediction with experiment is to let the two drive each other without a human in every loop. A self-driving laboratory couples a model that proposes the next experiment, usually through Bayesian optimisation over the design space, with robotics that prepare and test the proposed formulation and feed the result back to update the model. Each cycle sharpens the model where its uncertainty is highest, so the system converges on an optimum with far fewer experiments than a factorial design would require. Applied to nanomedicine, this closed-loop paradigm has been argued to compress formulation development that would traditionally take months into a fraction of the time, by concentrating experimental effort exactly where it is informative rather than spreading it evenly across a grid.¹⁰ The autonomous laboratory is the most complete expression of the data-driven vision, since it treats the model and the bench as a single instrument, and it is also the least widely deployed, for reasons that are organisational as much as technical.

3.5 Artificial intelligence in manufacturing and the digital thread

Formulation does not end at a validated recipe; it must be manufactured reproducibly at scale, and this is where machine learning meets the factory.



Continuous manufacturing, in which material flows through the process rather than moving in discrete batches, generates dense streams of process analytical data that are well suited to data-driven monitoring and control. Models trained on these streams can predict a critical quality attribute in real time from upstream sensor readings, flag a drift before it becomes a deviation, and in advanced implementations adjust process parameters to hold quality within specification. The digital twin, a calibrated computational replica of the process that runs

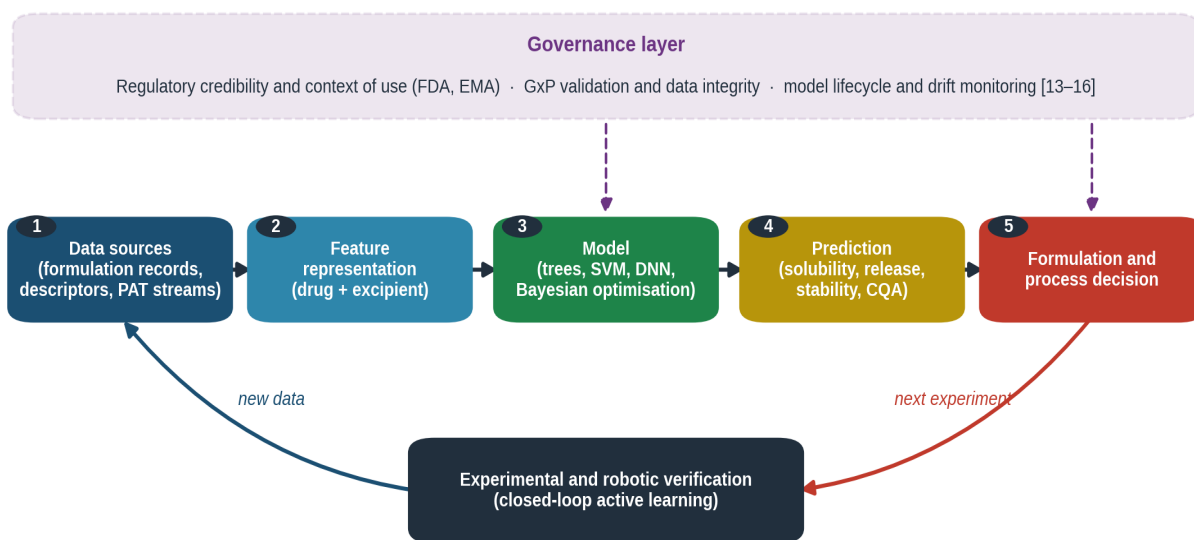
alongside the physical line, extends this into prediction and what-if analysis across the product lifecycle, and reviews now trace its role from development through to continuous production.¹¹ The manufacturing application matters to the adoption argument because it is where artificial intelligence touches the most heavily regulated part of the enterprise, and therefore where the gap between demonstration and deployment is both widest and most consequential.

Table 1. Representative applications of artificial intelligence across pharmaceutical formulation, showing the type of system, the approximate scale of the dataset used, the modelling method, and the principal reported outcome.

Application area	Representative system	Data scale	Method	Principal reported outcome	Ref
Solubility prediction	Unified web platform module	>4000 records across 27 solvents	Gradient-boosted trees, random forest	Test R^2 approx. 0.91	[6]
Cyclodextrin complexation	Unified web platform module	approx. 2990 records	Tree ensemble / kernel methods	Free-energy prediction, test R^2 approx. 0.87	[6]
In vitro formulation behaviour	Deep neural network	Paired drug and excipient descriptors	Deep learning	Formulation outcomes learnable from composition	[9]
Tablet disintegration	Sparse Bayesian regression	approx. 2000 data points	Sparse Bayesian learning with optimisation	Test R^2 0.96, MAPE approx. 12%	[8]
Solid dispersion stability	Stability classifier	Curated stability dataset	Ensemble classification	Stable vs unstable screening before stability studies	[7]
Autonomous nanomedicine	Self-driving laboratory	Grows online per cycle	Bayesian optimisation plus robotics	Optimum reached with far fewer experiments	[10]



Application area	Representative system	Data scale	Method	Principal reported outcome	Ref
Continuous manufacturing	Digital twin and process control	Real-time PAT data streams	Data-driven monitoring and control	Real-time CQA prediction and drift detection	[11]



Forward prediction narrows the design space; the closed loop couples model and bench so each experiment is spent where uncertainty is highest.

Figure 1. The data-to-decision workflow for artificial intelligence in pharmaceutical formulation, and the governance layer that constrains it.

4. The challenges that constrain model performance

Before the adoption barriers, which sit outside the model, it is worth being precise about the technical limits that sit inside it, because overstating what the models can do weakens the case for the parts that genuinely work.

4.1 The small-data problem

The single most important technical fact about formulation machine learning is that its datasets are small. Discovery models draw on chemical libraries of millions of compounds; a formulation dataset of a few thousand records is considered

large, and many useful ones number in the hundreds.⁶ The reason is intrinsic to the domain. Each data point is a physical experiment consuming material, instrument time and expert labour, and the most informative experiments, such as long-term stability studies, are the slowest and costliest to run. Small datasets punish complex models, which have the capacity to memorise the training set without learning a generalisable relationship, and this is precisely why the platforms that work best in this domain have favoured tree ensembles and kernel methods over deep networks.⁶ The field has, in



effect, discovered empirically that model capacity should be matched to data scale, and that the deep learning that transformed image and language tasks is often the wrong tool for a few hundred formulation records.

Several strategies mitigate the constraint without removing it. Careful feature engineering, in which molecular descriptors and physicochemical parameters encode prior knowledge so the model has less to learn from scratch, is standard practice and is part of why the better models succeed with modest data. Transfer learning, data augmentation and active learning offer further leverage, and the autonomous laboratory can be read as active learning made physical, since it spends each new experiment where the model is most uncertain.¹⁰ None of these makes small data behave like big data. They make small data go further, and the distinction matters for setting realistic expectations.

4.2 Generalisation, distribution shift and validation

A model that performs well on a held-out test set drawn from the same source as its training data has demonstrated interpolation, not necessarily the extrapolation that deployment demands. Formulation models are typically trained on the chemistries, excipients and processes represented in their datasets, and their reliability degrades, often silently, when they are applied to a new molecule class, a novel excipient or a different manufacturing route. This distribution shift is the practical reason a published accuracy figure cannot be taken at face value by a formulator

considering a genuinely new product. The honest reporting of applicability domain, the region of input space where a model's predictions can be trusted, is still uneven across the literature, and prospective validation on newly generated data, as opposed to retrospective testing on a held-out split, remains comparatively rare. The gap between retrospective and prospective performance is not a detail. It is the difference between a model that looks ready and a model that is ready.

4.3 The black box and the demand for explanation

Many high-performing models are opaque, returning a prediction without an interpretable account of how the inputs produced it. In a discovery setting this is often tolerable, since a wrong prediction costs a wasted synthesis. In formulation, and still more in manufacturing, an unexplained prediction that influences product quality is harder to accept, both because a formulator reasonably wants mechanistic insight before trusting a counterintuitive recommendation and because the regulatory apparatus expects decisions to be justifiable. Explainable artificial intelligence, the set of methods that attribute a prediction to its contributing features, has therefore become central rather than cosmetic, and its role as a bridge between predictive power and mechanistic understanding is an active field.¹² Feature-attribution analyses that identify, for instance, wetting time as the driver of disintegration show that explanation and accuracy need not trade off.⁸ The broader point is



that in a regulated domain, a model's interpretability is not a luxury added after the fact but a property that determines whether the model can be used at all.

5. Barriers to industrial adoption

The preceding challenges are real, but they are being addressed steadily by the research community and none of them alone explains why validated formulation models remain largely absent from routine development. The explanation lies in three barriers that sit outside the model and that the technical literature tends to underweight. These are the questions a pharmaceutical company must answer not to build a model but to use one.

5.1 Regulatory uncertainty and an evolving framework

For most of the period in which formulation machine learning matured, developers faced a regulatory vacuum: no agency had said clearly how an artificial intelligence model that influences a product decision would be evaluated. That vacuum is now filling, but with frameworks still in draft, which creates its own uncertainty. In 2023 the United States Food and Drug Administration issued a discussion paper on artificial intelligence in drug manufacturing that did not set rules but asked how the existing risk-based regulatory framework should apply to AI-enabled processes, and identified the ambiguities that continuous learning models and advanced process control create for a system built around fixed, validated procedures.¹³ Early in 2025 the agency followed with draft guidance

on the use of artificial intelligence to support regulatory decision-making for drug and biological products, organised around the concept of context of use and a risk-based assessment of a model's credibility for the specific question it is asked to answer.¹⁴ The European Medicines Agency reached a broadly compatible position in its 2024 reflection paper on artificial intelligence across the medicinal product lifecycle, which locates manufacturing and quality-control applications within the existing quality risk management principles of the relevant harmonised guidelines, calibrates scrutiny to patient risk and regulatory impact, and insists on a human-centric approach with traceable data governance throughout the model lifecycle.¹⁵

The convergence of these texts on a risk-based, context-of-use, human-oversight model is genuinely helpful, because it tells developers what will be asked of them. What it also reveals is how far most published models sit from being submittable. A model reported with a test-set accuracy and no defined context of use, no lifecycle plan, no applicability domain and no traceable data provenance is a research result, not a regulatory dossier. The barrier here is not that regulators have forbidden the technology. It is that meeting the emerging expectations requires evidence and documentation of a kind that academic development rarely generates, and that industrial adopters must now learn to produce.



5.2 Data integrity, validation and lifecycle management under the quality system

Even where a regulatory pathway is clear in principle, an artificial intelligence system must live inside a pharmaceutical quality system that was designed for deterministic, validated software and physical processes. This is a deeper obstacle than it first appears. Computerised systems in a regulated environment are validated to established frameworks, and the current good-practice guidance now explicitly addresses the validation of machine learning within a risk-based approach.¹⁶ A conventional system is validated once against a fixed specification and then locked. A machine learning model challenges that logic on several fronts at once. Its behaviour is defined by data rather than by explicit rules, so the training data itself becomes part of what must be controlled, documented and made traceable to the standards of data integrity that govern regulated records. A model that continues to learn after deployment is, from a validation standpoint, a system whose specification changes on its own, which the traditional validate-then-freeze paradigm cannot straightforwardly accommodate, and which is why the regulatory texts converge on keeping high-impact models frozen and monitored rather than continuously adapting.¹⁵ Model performance can drift as incoming material or process conditions move away from the training distribution, so a validated model requires ongoing monitoring of a kind that a validated pump does not. Each of these is soluble, but collectively they demand a validation and quality-assurance capability that most

organisations are still building, and the effort of building it is a real and often underestimated cost of adoption.

5.3 Organisational, talent and cultural barriers

The least technical barrier may be the most decisive. Deploying artificial intelligence in formulation is an organisational change, not merely a software installation, and organisational change is where many efforts stall. Analyses of the wider life sciences sector describe a consistent pattern in which numerous pilots demonstrate value but few scale into production, with the failure concentrated not in the modelling but in data readiness, integration with legacy systems, and the organisational capability to operationalise a model once it works.¹⁷ Formulation is a particularly demanding case, because success requires people who understand both the pharmaceuticals and the data science well enough to know when a model's recommendation is physically reasonable, and such people are scarce. A formulator without data fluency cannot judge a model's limits; a data scientist without formulation knowledge cannot encode the domain constraints that make a model trustworthy. The interface between the two disciplines is where trust is built or lost, and staffing that interface is a workforce problem that no algorithm solves.

Culture compounds the difficulty. Formulation development is a conservative discipline for sound reasons, since the cost of a quality failure is measured in patient harm and product recalls, and a field that has succeeded for decades with design of experiments has a rational reluctance to



hand decisions to a model it cannot fully interrogate. Adoption therefore depends on demonstrated reliability and on models that explain themselves, which loops back to the interpretability discussed above, and on leadership willing to invest in data infrastructure whose payoff is real but deferred. The proprietary and fragmented state of formulation data is itself partly an organisational artefact, since the most valuable datasets sit inside individual companies, are rarely shared for competitive reasons, and are often recorded in inconsistent formats that resist pooling even internally. The scarcity of large, standardised, openly available formulation datasets is the data problem of section 4.1 seen from the industry side, and it is sustained by organisational incentives rather than by any technical impossibility.

5.4 Economic, intellectual-property and integration costs

Adoption also carries costs and uncertainties that are financial and legal rather than scientific, and these shape decisions at the level where budgets are set. Building the capability is expensive before it is productive, because it requires data

infrastructure, computational tooling, validated systems and specialist staff whose combined outlay precedes any measurable return, and the return itself is a reduction in development risk and time that is real but difficult to attribute cleanly to the model. This deferred and diffuse payoff competes for investment against projects with nearer and more legible returns, which is one reason pilots are easier to fund than the infrastructure that would let a pilot scale. Intellectual property adds a further layer of uncertainty, since the status of a formulation conceived by a generative model, the ownership of models trained on pooled or third-party data, and the confidentiality of proprietary formulation data used to train shared systems are all questions the field has not fully resolved, and unresolved questions of this kind make legal and commercial functions cautious. Integration with entrenched laboratory and manufacturing information systems is a mundane but persistent cost, because a model delivers value only when its inputs and outputs flow through the systems that formulators and operators actually use, and retrofitting that connectivity onto legacy infrastructure is frequently the step where a promising pilot quietly stalls.

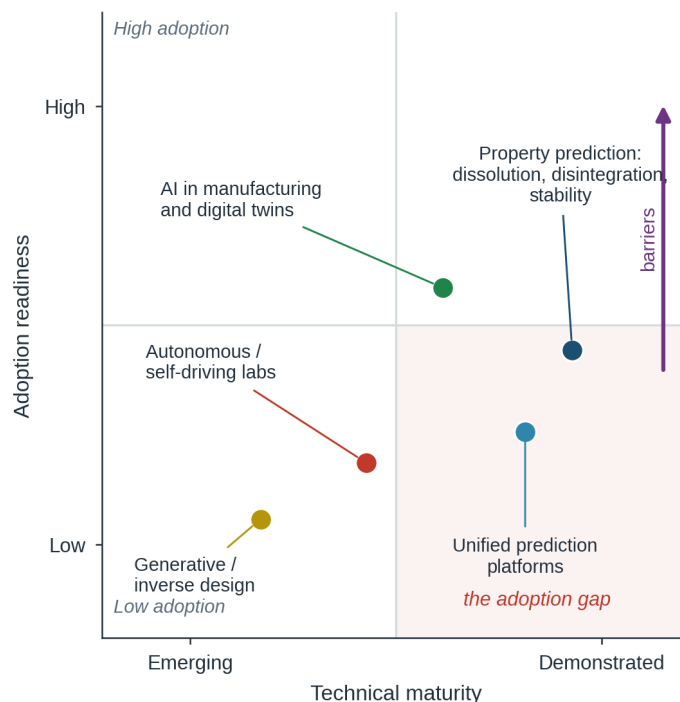


Figure 2. The adoption gap between technical maturity and adoption readiness across artificial intelligence applications in formulation.

6. Discussion

Read together, the evidence supports a clear judgement. The methodological case for artificial intelligence in pharmaceutical formulation is strong and no longer seriously in doubt at the level of demonstration. Models predict formulation-relevant properties with accuracies that are useful in practice, unified platforms show that these capabilities can be integrated, and autonomous laboratories show that prediction and experiment can be coupled into a system that finds good formulations with markedly less experimental effort.^{6,8,10} A formulator who ignores these tools increasingly leaves efficiency on the table.

The weight of evidence equally supports a second judgement that the field's promotional literature tends to soften. The binding constraints on adoption are downstream of the model. The

small-data reality caps how far the methods generalise and rewards discipline over sophistication in model choice.⁶ The regulatory frameworks now emerging ask for lifecycle evidence, defined context of use and traceable data governance that most published models do not provide.¹³⁻¹⁵ The quality-system requirements of validation, data integrity and drift monitoring impose obligations that a research prototype never has to meet.¹⁶ And the organisational realities of talent, culture, legacy integration and data fragmentation determine, more than any accuracy figure, whether a working model becomes a working practice.¹⁷ These are not reasons for pessimism. They are a correct diagnosis of where the remaining work lies, and it lies mostly in infrastructure, process, incentives and people rather than in algorithms.



Two weaknesses in the current literature deserve emphasis because they distort the field's self-assessment. The first is the scarcity of prospective validation. A great many models are reported on retrospective held-out splits, and comparatively few are tested on formulations designed and made after the model was fixed, which is the only test that mirrors deployment. The second is the absence of shared benchmarks. Because datasets are proprietary and non-standardised, published accuracies are computed on incommensurable data and cannot be compared across studies, which makes it genuinely difficult to know which methods are best and easy for the field to mistake activity for progress. Both weaknesses are consequences of the data situation, and both would be substantially relieved by the same remedy.

7. Future directions

The most valuable single intervention would be shared, standardised, openly accessible formulation datasets, built to consistent reporting standards and, where possible, through precompetitive collaboration among companies that individually hold fragments of the picture. Such datasets would relieve the small-data constraint, enable the common benchmarks the field lacks, and make prospective validation studies feasible at scale. Standardised reporting of what a model was trained on, where its applicability domain lies, and how it performs prospectively should become an expectation of publication rather than an exception, so that a reported result carries the information a potential adopter actually needs. Regulatory engagement

should be proactive rather than reactive, using the mechanisms that agencies are opening to test AI-assisted formulation workflows against the emerging context-of-use and credibility frameworks before a product filing depends on them, so that the documentation burden is understood early.^{14,15} Within companies, the interface between formulation and data science should be staffed and trained deliberately rather than left to chance, and validation and quality functions should develop the specific competence to assess machine learning systems, since these capabilities are prerequisites for adoption rather than by-products of it. Finally, the field should invest in interpretable and hybrid models that combine data-driven prediction with mechanistic and physical knowledge, because such models generalise better from small data, are easier to validate, and are more likely to earn the trust of both formulators and regulators.¹² None of these is a research problem in the conventional sense. Each is a matter of building the conditions under which the research results already in hand can be used.

8. Conclusion

Artificial intelligence has demonstrably become a real and capable method in pharmaceutical formulation, able to predict product performance, accelerate design and, in its most advanced form, run the design loop autonomously. The reason it has not yet transformed routine industrial practice is not that the models fail but that the surrounding conditions for their use are still being built. The decisive barriers are the small and fragmented state of formulation data, the



evolving and demanding regulatory expectations, the requirements of validation and data integrity within the quality system, and the organisational capability to operationalise a model at scale. The single most important takeaway is that the frontier of this field has shifted from the algorithm to everything around it, and progress now depends less on better models than on better data, clearer regulatory pathways, and the people and processes to put working models to work.

Declarations

Conflicts of interest: None declared.

Funding: No specific funding was received for this work.

References

1. Paul SM, Mytelka DS, Dunwiddie CT, Persinger CC, Munos BH, Lindborg SR, Schacht AL. How to improve R&D productivity: the pharmaceutical industry's grand challenge. *Nat Rev Drug Discov.* 2010;9(3):203-214. doi:10.1038/nrd3078.
2. Vamathevan J, Clark D, Czodrowski P, et al. Applications of machine learning in drug discovery and development. *Nat Rev Drug Discov.* 2019;18(6):463-477. doi:10.1038/s41573-019-0024-5.
3. Mak KK, Pichika MR. Artificial intelligence in drug development: present status and future prospects. *Drug Discov Today.* 2019;24(3):773-780. doi:10.1016/j.drudis.2018.11.014.
4. Vora LK, Gholap AD, Jetha K, Thakur RRS, Solanki HK, Chavda VP. Artificial intelligence in pharmaceutical technology and drug delivery design. *Pharmaceutics.* 2023;15(7):1916. doi:10.3390/pharmaceutics15071916.
5. Warke S, Katari O, Jain S. Current status on the convergence of artificial intelligence and formulation development in industry: a review. *AAPS PharmSciTech.* 2025;26:44. doi:10.1208/s12249-025-03296-0.
6. Dong J, Wu Z, Xu H, Ouyang D. FormulationAI: a novel web-based platform for drug formulation design driven by artificial intelligence. *Brief Bioinform.* 2024;25(1):bbad419. doi:10.1093/bib/bbad419.
7. Han R, Xiong H, Ye Z, Yang Y, Huang T, Jing Q, et al. Predicting physical stability of solid dispersions by machine learning techniques. *J Control Release.* 2019;311-312:16-25. doi:10.1016/j.jconrel.2019.08.030.
8. Ghazwani M, Hani U. Prediction of tablet disintegration time based on formulations properties via artificial intelligence by comparing machine learning models and validation. *Sci Rep.* 2025;15:13789. doi:10.1038/s41598-025-98783-6.
9. Yang Y, Ye Z, Su Y, Zhao Q, Li X, Ouyang D. Deep learning for in vitro prediction of pharmaceutical



- formulations. *Acta Pharm Sin B*. 2019;9(1):177-185.
doi:10.1016/j.apsb.2018.09.010.
10. Hickman RJ, Bannigan P, Bao Z, Aspuru-Guzik A, Allen C. Self-driving laboratories: a paradigm shift in nanomedicine development. *Matter*. 2023;6(4):1071-1081.
doi:10.1016/j.matt.2023.02.007.
11. Maharjan R, Kim NA, Kim KH, Jeong SH. Transformative roles of digital twins from drug discovery to continuous manufacturing: pharmaceutical and biopharmaceutical perspectives. *Int J Pharm X*. 2025;10:100409.
doi:10.1016/j.ijpx.2025.100409.
12. Lavecchia A. Explainable artificial intelligence in drug discovery: bridging predictive power and mechanistic insight. *WIREs Comput Mol Sci*. 2025;15(5):e70049.
doi:10.1002/wcms.70049.
13. US Food and Drug Administration. Artificial intelligence in drug manufacturing: discussion paper. Docket FDA-2023-N-0487. Silver Spring (MD): FDA; 2023.
14. US Food and Drug Administration. Considerations for the use of artificial intelligence to support regulatory decision-making for drug and biological products: draft guidance for industry. Silver Spring (MD): FDA; 2025.
15. European Medicines Agency. Reflection paper on the use of artificial intelligence (AI) in the medicinal product lifecycle. EMA/CHMP/CVMP/83833/2023. Amsterdam: EMA; 2024.
16. International Society for Pharmaceutical Engineering. GAMP 5: a risk-based approach to compliant GxP computerized systems. 2nd ed. North Bethesda (MD): ISPE; 2022.
17. Chui M, Roberts R, Yee L, et al. Scaling gen AI in the life sciences industry. McKinsey & Company; 2024.