



Journal of Advanced Pharmaceutical Sciences and Natural Products (JAPSNP)

CLOSED-LOOP FORMULATION DEVELOPMENT THROUGH ARTIFICIAL INTELLIGENCE: BAYESIAN OPTIMIZATION, ROBOTICS, AND SELF-DRIVING LABORATORIES IN PHARMACEUTICS

Dr. Vishal Chauhan*

MAHSA University, Jln SP 2, Bandar Saujana Putra, 42610, Jenjarom, Selangor, Malaysia

Corresponding author:

Dr. Vishal Chauhan

Email: vishal@mahsa.edu.my

ABSTRACT

Closed-loop experimentation replaces the human step in the formulation cycle with an algorithm that selects the next composition and hardware that prepares and measures it without further instruction. This review assesses what that idea has delivered in pharmaceuticals, and separates the three capabilities a working loop requires: a decision policy that is sample-efficient under uncertainty, execution hardware that can act without a person in the path, and a data layer that makes every outcome, including every failure, machine-readable. Evidence is now solid for the first. Bayesian optimization reached optimal orally disintegrating tablet conditions in roughly ten experiments where a factorial design needed about twenty-five, produced a monoclonal antibody formulation optimized simultaneously for thermal stability, colloidal interaction and

interfacial stress in thirty-three experiments, and located high-solubility injectable vehicles after sampling 256 of 7,776 combinations. A robotic tableting platform has carried six active ingredients from raw material characterisation to in-specification tablets within six hours while cutting material use by 65%. The second capability is advancing but remains limited by the physical awkwardness of pharmaceutical matter, particularly cohesive powders and low-dose solids. The third is largely absent. The central argument is that the binding constraint is metrological rather than computational: dissolution behaviour, physical stability over months and performance in a patient cannot be measured on the timescale a loop consumes, so every autonomous campaign optimizes a surrogate. Managing the distance between surrogate and endpoint, together with reporting sample efficiency against stated baselines, will decide whether these systems produce better medicines or only faster experiments.

Keywords: drug formulation; Bayesian optimization; machine learning; laboratory automation; robotics; quality by design; pharmaceutical technology.



1. Introduction

Formulation sits in an awkward place in drug development. It is downstream of the decisions that attract investment and upstream of the trials that generate evidence, yet it determines whether a promising compound can be given to a patient at all. Between 70% and 90% of new chemical entities entering development are poorly soluble, and pharmaceutical research spending has risen roughly 51-fold over four decades while clinical success rates have stayed near 10%¹. The capitalised cost of bringing one medicine to market has been estimated at a median of about 1.1 billion US dollars across 63 agents approved between 2009 and 2018². Development economics are dominated by attrition, and attrition is partly a formulation problem.

The methodological answer pharmaceuticals adopted was quality by design. Under ICH Q8(R2) a formulator identifies critical material attributes and process parameters, explores them through a structured experimental design, and defines a design space, described in the guideline as the multidimensional combination and interaction of input variables demonstrated to provide assurance of quality³. The approach works and is regulatorily well understood. Its limits are equally clear. A factorial or response-surface design assumes a low-order polynomial describes the response surface, that the factor count stays small enough for the design to remain affordable, and that the whole design can be specified before any result is seen. Nanomedicines, biologics and multifunctional excipient systems violate all three⁴.

An adaptive alternative has existed since the late 1990s. Efficient global optimization builds a probabilistic surrogate of an expensive function and uses its own uncertainty to choose the next sample, converging with a fraction of the evaluations a fixed design needs⁵. Modern Bayesian optimization descends from that idea⁶. Its move into the wet laboratory was slow, but by 2021 it had beaten trained chemists at reaction optimization in a controlled comparison⁷. The pharmaceutical demonstration came earlier and drew less attention: applied to the formulation and process parameters of an orally disintegrating tablet, Bayesian optimization reached the optimum in about ten experiments where design of experiments required about twenty-five⁸.

What turns an optimization algorithm into a closed loop is the removal of the human from the execution path. A self-driving laboratory couples the decision policy to hardware that prepares samples, runs analyses and returns results directly to the algorithm, completing a design, make, test and analyse cycle without instruction⁹. The concept matured in materials chemistry: a mobile robot working with standard laboratory instruments ran 688 photocatalysis experiments across eight days and a ten-variable space, improving activity roughly six-fold¹⁰, and reviews of the area report acceleration between ten and a thousand fold alongside build costs from under a thousand dollars to over a million¹¹.

Pharmaceuticals arrived late, and the delay has a technical cause rather than only a cultural one. A



photocatalyst is judged by one measurable rate. A tablet is judged by content uniformity, hardness, friability, disintegration, dissolution, stability across months of storage, manufacturability at scale and performance in a patient. Some of those can be measured in seconds and some cannot be measured before the loop must make its next decision.

There is also a reporting problem worth naming at the outset. Much of the enthusiasm for autonomous experimentation rests on acceleration factors quoted without a stated denominator, and a tenfold speedup against an unspecified baseline is not a measurement. Where a comparator is reported, the advantage is real but more modest than the headline figures suggest, and where none is reported the claim cannot be assessed at all. A review that repeats those factors uncritically does the field no service.

This review therefore does three things. It decomposes the loop into capabilities that are separable and separately mature. It assesses what Bayesian optimization and robotic execution have actually delivered in formulation, with attention to the baselines behind claimed speedups. And it argues that the binding constraint is measurement rather than computation, which makes the design of fast, loop-compatible surrogate assays the research programme that matters most. Scope is drawn around drug product development; molecular design and clinical trial optimization are excluded, as is property prediction where no

experimental loop is closed, except where such models serve as the surrogate inside a loop.

2. Anatomy of the loop, and levels of autonomy

A closed-loop campaign is usually drawn as a circle with four arrows. That picture hides the components where campaigns fail. Figure 1 gives a more useful decomposition into eight parts: the objective, which converts a target product profile into something an algorithm can maximise; the parameterisation of the design space; the surrogate model, which predicts the response and its own uncertainty; the acquisition function, which turns prediction and uncertainty into a decision; the execution hardware; the analytics; the data layer, which stores proposals, executed conditions and results including failures; and the stopping rule, which in practice is usually the exhaustion of a budget rather than any statistical criterion.

Two of these deserve more attention than they usually receive. Defining the objective is not clerical work. A formulator who wants a stable, redispersible, injectable suspension is expressing several objectives with implicit priorities and hard constraints, and the translation into a scalar or vector quantity determines what the loop will actually pursue rather than merely how it will pursue it. The stopping rule is the other. A campaign that halts when its budget runs out has not established that further sampling would be unproductive, and a campaign that halts when improvement stalls may simply have found a local optimum that its own uncertainty model



cannot see past. Neither convention is wrong, but which one was used should be reported, and it usually is not.

Two feedback paths matter. The fast path returns an analytical result to the surrogate within one cycle. The slow path returns consolidated

knowledge to the data layer, where it becomes a prior for later campaigns. Nearly all published pharmaceutical work exercises the first. Almost none exercises the second, which is why each new campaign starts from something close to ignorance.

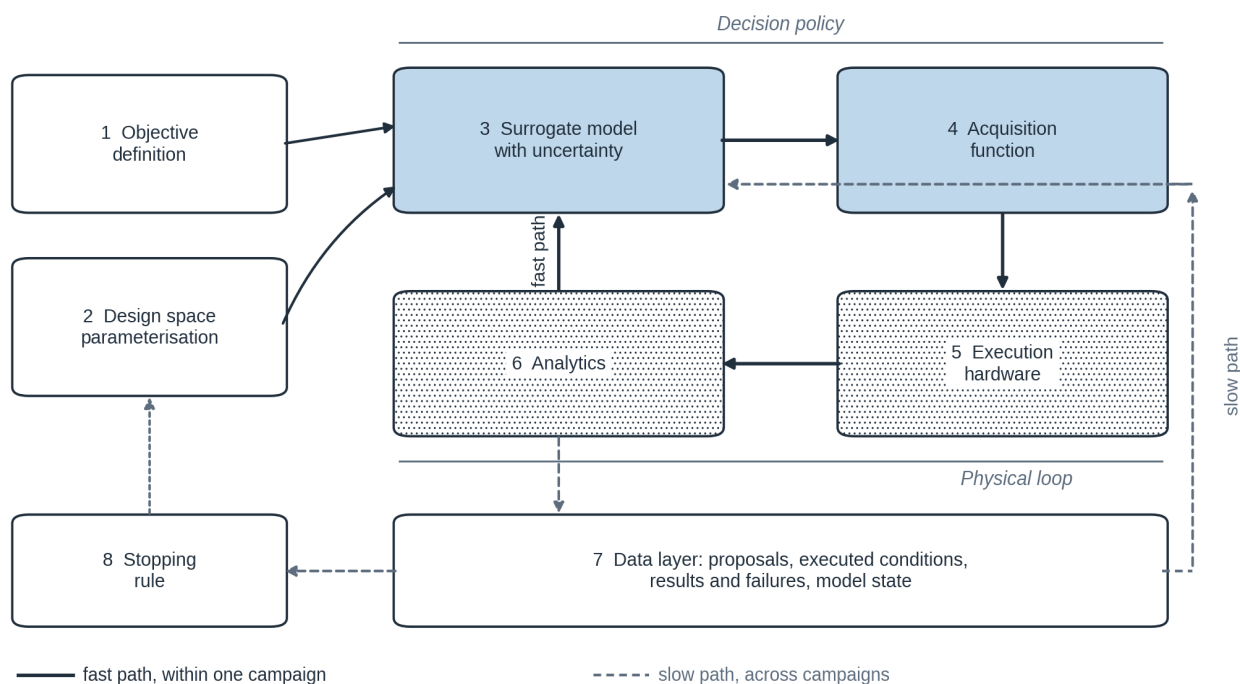


Figure 1. The eight components of a closed-loop formulation campaign, with the fast and slow feedback paths distinguished. The objective definition converts a target product profile into a scalar or vector quantity; the design space parameterisation fixes which variables are continuous, discrete or categorical; the surrogate model and acquisition function together constitute the decision policy; execution hardware and analytics constitute the physical loop; the data layer stores proposals, executed conditions and outcomes including failures; and the stopping rule terminates the campaign. The fast path (solid arrows) returns an analytical result to the surrogate within one cycle and is exercised by essentially all published pharmaceutical campaigns. The slow path (dashed arrows) returns consolidated knowledge to the data layer so that it becomes a prior for later campaigns, and is exercised by almost none.

The term self-driving laboratory is applied to systems that differ by orders of magnitude in what they do unattended, so a scale is needed. A published taxonomy adapts the automotive convention: level 1 is assisted operation such as a liquid handler, level 2 is partial autonomy such as machine-generated protocols, level 3 is

conditional autonomy in which the system completes at least one full cycle of the scientific method and calls for help only on anomalies, level 4 automates protocol generation and hypothesis revision, and level 5, an autonomous researcher, has not been achieved¹². Level 3 is



the stated minimum for a system to count as a self-driving laboratory.

This distinction exposes how much published formulation work sits below the threshold. A study that trains a model on literature data, proposes candidates, and has a graduate student prepare them is at level 1 or 2 however the abstract is worded. Figure 3 places the pharmaceutical evidence on that scale, and the crowding at the lower levels is the finding.

Every loop optimizes something measurable, and in pharmaceuticals that quantity is almost never the property of clinical interest. Particle size stands in for biodistribution, a dissolution profile for absorption, a melting temperature for two-year stability. A formulator applies judgement about when a surrogate has stopped tracking the endpoint. An algorithm applies none, and will exploit any region where the two diverge. Section 5 returns to this, because it is the argument on which the review turns.

3. The decision engine

Bayesian optimization is built for a specific situation: the objective is expensive to evaluate, has no analytical gradient, is observed with noise, and the evaluation budget is small. Formulation matches that description closely. One tablet formulation consumes active ingredient, occupies a press and a dissolution bath, and yields a single noisy observation of several correlated responses. Thirty to a hundred experiments is a realistic budget; ten thousand is not.

The machinery is simple in outline. A Gaussian process places a prior over functions, and after each observation the posterior gives a predicted response and a calibrated uncertainty at every untried point. An acquisition function turns these into a score, and the next experiment is the point that maximises it ^{5,6}. Which acquisition function is chosen matters far less in practice than three structural features of formulation problems, all of which are present at once.

Multi-objective structure is intrinsic. Raising drug loading in a nanoparticle typically raises size; raising pH to improve thermal stability may worsen colloidal interaction. Collapsing several objectives into one weighted score forces the formulator to fix trade-offs before the shape of the trade-off surface is known. Approximating a Pareto front instead makes the conflict visible, as in a monoclonal antibody campaign where pH moved melting temperature and the diffusion interaction parameter in opposite directions ¹³.

Constraints are of two kinds. Some are known and cheap to check, such as excipient fractions summing to unity or an osmolality window. Others are unknown until an experiment is attempted, such as the composition at which a suspension gels or an extrudate fails to form. Placing a Gaussian process prior on the constraint alongside the objective allows feasibility to be learned rather than assumed ¹⁴. The harder case is unresolved: practical analysis of self-driving laboratory operation notes that optimizers generally assume a search space without holes and with fully known constraints, whereas real campaigns meet infeasible regions



created by reaction scope, measurement failure and reagent availability ¹⁵.

Mixed variables are where formulation departs furthest from the problems these algorithms were designed on. Choosing between three polymers is not a continuous decision, and a numerical index over polymer identity implies an ordering that does not exist. Methods that treat categorical variables natively, optionally informed by physicochemical descriptors so that similar excipients share information, are a requirement rather than a refinement ¹⁶. Software now unifies mixed parameters with multi-objective, constrained and multi-fidelity operation in one package, and supports warm starting a campaign from related historical data ¹⁷.

A fourth feature is less often discussed and matters more in industry than in academia. The sample efficiency of a campaign depends heavily on where it starts. A Gaussian process with an uninformative prior spends its first several experiments learning what an experienced formulator already knows, and in a budget of thirty experiments those are expensive lessons. Two routes exist for supplying that knowledge: encoding expert expectations about promising regions directly into the search, or warming the surrogate with data from related campaigns, which is the meta-learning capability now built into orchestration libraries ¹⁷. The pharmaceutical industry is unusually well placed to exploit the second route and unusually poor at doing so. Decades of chemistry, manufacturing and controls development have produced formulation and stability data on tens of

thousands of compositions, held in study reports and regulatory dossiers in a form no algorithm can read. Extracting even a fraction of that into a machine-readable prior would improve sample efficiency more than any change of acquisition function.

3.1 What the sample-efficiency evidence shows

Claims of acceleration are the currency of this literature and deserve scrutiny, because the baseline is frequently unstated. Three classes of comparison appear and they are not equivalent, as Figure 2 shows.

The strongest is a head-to-head against a conventional design on the same problem. The orally disintegrating tablet study is of this type, reporting roughly ten experiments against about twenty-five for design of experiments, and showing that periodic re-tuning of the surrogate's hyperparameters during a campaign stabilises performance across acquisition function and training set choices ⁸.

The second reports the fraction of a defined space explored, with a control. A closed-loop campaign for metallophotocatalyst formulations synthesised 55 of 560 candidate molecules across seven batches, raising yield from 39% to 67%, then evaluated 107 of 4,500 possible condition combinations to reach 88% against a random-sampling control that reached only 75% ¹⁸. A semi-automated formulator for injectable vehicles sampled 256 of 7,776 combinations, about 3%, returning seven leads above 10 mg/mL solubility within days ¹⁹. These are informative because the denominator is stated.



The third reports elapsed time or an absolute experiment count with no comparator. A biologics campaign reaching a three-objective optimum in 33 experiments is impressive on its face¹³, and a robotic platform identifying nine

acceptable recipes in 15 working days is a real result²⁰, but neither states what a conventional programme would have consumed, so the acceleration factor is an inference rather than a measurement.

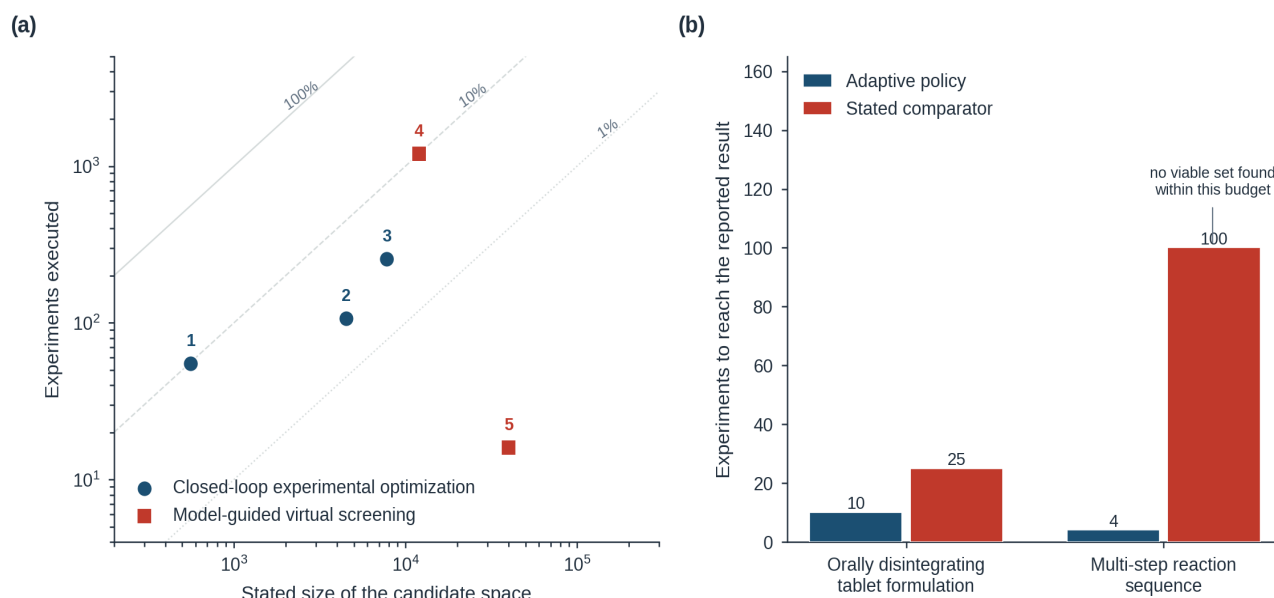


Figure 2. Reported experimental budgets for optimization campaigns in formulation and closely related chemistry, showing how rarely a comparator is given. (a) Experiments executed against the stated size of the candidate space, on logarithmic axes, for the campaigns reporting both; diagonal guides mark 100%, 10% and 1% of the space evaluated. Numbered points are 1 and 2, the two stages of a closed-loop catalyst formulation campaign¹⁸; 3, injectable solubility vehicles¹⁹; 4, robotic lipid library synthesis ranked against a virtual library²¹; and 5, model-guided ionizable lipid discovery²². Circles mark campaigns that closed an experimental loop, squares those that screened a virtual library and synthesised only the top-ranked candidates. (b) The two campaigns in this set reporting a direct comparison on the same problem, expressed as experiments consumed by the adaptive policy against its stated comparator: Bayesian optimization against design of experiments for orally disintegrating tablets⁸, and reinforcement learning against Bayesian optimization on a multi-step reaction space where the latter found no viable condition set within its budget²³.

Two cautions follow. Benchmarks built from synthetic test functions overstate performance on real, noisy experimental surfaces. And the advantage of a sophisticated variant over a simple one is often problem-dependent: a systematic assessment of multi-fidelity Bayesian optimization across three real discovery problems concluded that its benefit depends on how cheap and how informative the low-fidelity source actually is, and is not universal²⁴.

3.2 Where the decision engine fails

Four failure modes recur. Gaussian process performance degrades as dimensionality grows, so practical campaigns depend on a screening step that reintroduces the human judgement the loop was meant to remove. Uncertainty misspecification is more insidious, since a poorly chosen kernel produces confident predictions where the model has no information, and the symptom is a campaign that converges quickly



onto a local optimum and reports success. Unknown infeasibility is the third¹⁵. The fourth is a mismatch between policy and problem: where the task is sequential and combinatorially structured, other policies do better, and a reinforcement learning agent navigating a multi-step reaction space of more than 40 dimensions found a viable condition set after four experiments and reached 94% of the known

optimum within 100 trials, on a problem where standard Bayesian optimization had found nothing viable in 100 experiments²³.

4. Execution hardware

Execution platforms fall into four archetypes, compared in Table 1, and the choice among them is dictated by the physical form of the material rather than by the sophistication of the algorithm.

Table 1. Execution architectures for closed-loop formulation, with the properties that determine suitability for pharmaceutical work. Cost figures are those reported in the cited sources and span published extremes rather than representing typical installations.

Architecture	Reported cost	Throughput characteristics	Suitability for pharmaceuticals	Evidence
Fixed workcell	Specialised commercial systems from around US\$1 million; conventional laboratory equipment around US\$800,000 with US\$80,000 annual maintenance ¹²	Highest for the single workflow it was built for	Strong for one dosage form, poor flexibility across forms	Autonomous tableting from raw material to specification in 6 h across six active ingredients ²⁵
Mobile robot operating conventional instruments	General-purpose arms from about US\$10,000; open hardware from US\$400 ¹²	688 experiments in 8 days for a ten-variable space ¹⁰	Strong, because existing qualified instruments are retained rather than replaced	Photocatalyst optimization ¹⁰ ; exploratory chemistry sharing instruments with human users ²⁶
Flow and microfluidic platform	Entry-level liquid handling from about US\$5,000 ²⁷	Nanoparticle synthesis platforms reported at more than 1,000 times conventional rates ¹¹	Strong for liquids, nanoparticles and continuous processes; unsuited to powders	Reviewed acceleration of 10 to 1,000 fold across self-driving platforms ¹¹
Cloud laboratory	Subscription from around US\$50,000 per month ¹²	Set by the provider; one comparison cites 46,620 against 8,880 samples annually ¹²	Immediate access without capital outlay, at a price that excludes most academic groups	Accessibility assessed as aspirational rather than achieved ²⁸

Automation literature from chemistry is dominated by liquids, and that conceals the problem pharmaceuticals actually has. Solid dosage forms depend on powders that bridge, adhere, segregate, absorb moisture and flow unpredictably at the milligram scale. Accurate low-mass solid dispensing has been named explicitly as an unresolved bottleneck limiting closed-loop reliability, alongside automated identification of unexpected byproducts¹⁵. Semi-solids present the same difficulty with viscosity,

and suspensions can change between preparation and measurement, so the loop's observation depends on when it looked.

There is also a maintenance problem that publications systematically understate. Between the algorithm and the instruments sits orchestration software that schedules, dispatches, handles exceptions and recovers from them, and this is the component most often improvised and least often described. A campaign is only



autonomous for as long as nothing unexpected happens, and in a laboratory something unexpected happens routinely: a tip box empties, a syringe clogs, an instrument returns a result outside the range its parser expects. Assessments of what limits practical deployment place these operational failures alongside the algorithmic ones^{15,28}. Any laboratory that has watched a robotic campaign stop overnight because a consumable ran out will recognise the diagnosis, and it explains why the gap between demonstrated and sustained throughput is wide.

Cost is quoted selectively in this field, usually at the top of the range. Bespoke platform

development is estimated at one to two years, which hardware standardisation could in principle reduce to one to two months¹¹, while entry-level liquid handling robots are available for about five thousand dollars²⁷. An assessment of the field's accessibility is candid that high development costs, limited open-source tooling, workforce training gaps and unresolved data standardisation leave broad access an aspiration rather than an achievement²⁸.

5. Evidence in pharmaceuticals

Table 2. Machine-learning-guided and closed-loop formulation studies in pharmaceuticals, with the decision policy, degree of physical automation, size of the space addressed and experimental budget as reported by each source. Autonomy level follows the five-level scale of Tobias and Wahab¹², where level 3 is the minimum qualifying as a self-driving laboratory. Entries marked prediction only did not close an experimental loop, and are included because they define the surrogate models on which future loops depend.

System	Decision policy	Physical automation	Space addressed	Budget consumed	Principal reported outcome	Autonomy	Ref.
ODT formulation and process	BO with periodic hyperparameter re-tuning	None reported	Formulation and process parameters	About 10 experiments	Optimum reached against about 25 for design of experiments	1 to 2	⁸
Tablets from six active ingredients	Physics-informed BO of excipient choice and compaction pressure	Full robotic tableting data factory	Nine formulation cases	4 to 6 calibration experiments per case	In-specification tablets in 6 h; material use cut by 65%; R ² 0.90 porosity, 0.89 tensile strength	3 to 4	²⁵
Injectable solubility vehicles	Sampling with human re-planning	Semi-automated robotic formulator	7,776 excipient combinations	256 sampled, about 3%	Seven leads above 10 mg/mL within days	2 to 3	¹⁹
Formulated liquid product	Classification model with multi-objective optimization	Semi-automated robotic platform	Continuous and discrete outputs, no physical model	15 working days	Nine recipes meeting customer criteria	2 to 3	²⁰



			available				
mAb formulation	Multi-objective BO under osmolality and pH constraints	High-throughput screening	Melting temperature, diffusion interaction parameter, interfacial stability	33 experiments	Optimized formulation with trade-offs quantified	2 to 3	¹³
Ionizable lipids for mRNA delivery	Deep learning ranking of a virtual library	Robotic liquid handling for one-pot synthesis	About 12,000 virtual candidates	1,200 lipids synthesised in one day; 15 validated per cell type	Lead gave 7.8-fold greater muscle delivery than comparator	2	²¹
Ionizable lipid discovery	Model-guided virtual screening	Combinatorial synthesis	40,000 virtual designs from a 584-lipid training library	Top 16 synthesised	Lead outperformed benchmarks for muscle and immune-cell transfection	1 to 2	²²
LNP in vivo screening	Barcoded pooled readout	Automated mass spectrometry readout	384 ionizable lipids, more than 400 formulations	Nine mice	High-throughput in vivo ranking without sequential testing	1	²⁹
Printed drug delivery systems	Supervised models including neural networks	None	968 formulations from 114 publications	Literature dataset	Up to 93% accuracy for processability; release time error about 24 min	Prediction only	³⁰
Long-acting injectables	Gradient boosting, prospectively validated	Conventional synthesis	43 drug-polymer combinations	181 release profiles, 3,783 measurements	Two prospective formulations matched predicted release	Prediction only	³¹

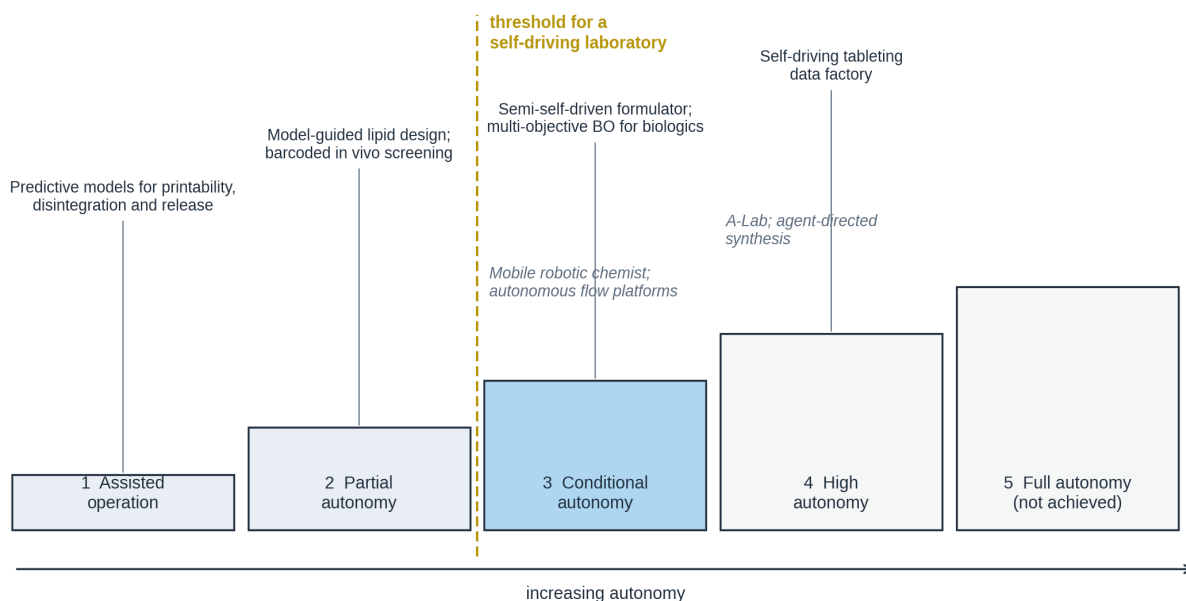


Figure 3. Published pharmaceutical formulation work placed on the five-level laboratory autonomy scale, where level 3, conditional autonomy, is the stated minimum for a system to qualify as a self-driving laboratory and level 5 has not been achieved in any field¹². Most formulation studies described as artificial-intelligence-driven occupy levels 1 and 2, where a model proposes candidates that people then prepare and interpret. The tableting data factory is the only pharmaceutical example in this review spanning the level 3 threshold across a complete dosage form workflow²⁵. Non-pharmaceutical exemplars from chemistry and materials science are shown in grey for calibration.

5.1 Nanomedicine

Nanomedicine is where formulation data science is most advanced, for a straightforward reason: size, polydispersity, zeta potential, encapsulation efficiency and in vitro transfection are fast, quantitative and automatable.

Two strategies have emerged and they solve different halves of the problem. The first uses machine learning to choose molecules for synthesis. A combinatorial library of 584 ionizable lipids trained a model that screened 40,000 virtual designs, of which the top 16 were made and tested²². A deep learning platform synthesised 1,200 ionizable lipids in a single day by one-pot chemistry with robotic liquid handling, ranked a virtual library of about

12,000, and validated fifteen candidates per cell type, the lead achieving 7.8-fold greater muscle delivery than a standard comparator²¹.

The second attacks the measurement bottleneck instead, by making the in vivo test high-throughput. Peptide-encoding messenger RNA barcodes read by mass spectrometry allowed a library of more than 400 formulations built from 384 ionizable lipids to be screened in nine mice²⁹. This work deserves more attention from the closed-loop community than it receives, because it addresses the constraint Section 6 identifies as binding.

The two strategies are complementary rather than competing, and the distinction between them is worth holding onto because they are often



described in the same terms. Model-guided library design compresses the search over chemical space, and its bottleneck is the quality of the training set. Pooled in vivo readout compresses the measurement, and its bottleneck is whether the barcode faithfully reports what a separately dosed formulation would have done. Neither, on its own, closes a loop: in both cases the model proposes and a person interprets, which is why these studies sit at levels 1 and 2 in Table 2 despite the sophistication of the modelling. A genuine nanomedicine loop would couple robotic formulation to a pooled readout with the ranking fed straight back into the next design round, and that combination has not yet been reported.

5.2 Oral solids, biologics and liquids

Oral solids concentrate the physical difficulties, and until recently the literature was dominated by prediction rather than experimentation. Models mining 968 formulations from 114 publications reached up to 93% accuracy for hot-melt extrusion processability, with drug release times predicted to a mean error of roughly 24 minutes³⁰.

The decisive change came with a platform that closes the loop on tablets themselves. A digital formulator combined with a self-driving tableting data factory carried nine formulation cases across six active ingredients from raw material characterisation to in-specification tablets in six hours, using Bayesian optimization of excipient selection and compaction pressure, cutting material consumption by 65% against current practice, with porosity and tensile

strength models reaching coefficients of determination of 0.90 and 0.89 from only four to six calibration experiments per case²⁵. This is the most complete demonstration of autonomous formulation of a conventional dosage form published to date, and its efficiency claim rests on material consumption rather than on an acceleration factor. That choice of metric is itself instructive. Material consumption is the constraint that actually binds in early development, where a few hundred milligrams of a candidate may be all that exists, and it can be measured unambiguously, whereas time saved depends entirely on what the comparator programme would have done.

What that platform closes and what it leaves open is worth stating precisely, because the distinction generalises. It closes the loop on properties measurable at the moment of manufacture: porosity, tensile strength, and whether a compact forms at all. It does not close the loop on dissolution across a full profile, on physical stability, or on content uniformity at scale, and it does not claim to. A formulation reaching specification in six hours is a candidate for development rather than a finished product, and the honest reading of the result is that the platform compresses the front half of the workflow while leaving the back half exactly as long as it was. That is still a substantial gain, since the front half is where most compositions are discarded.

Protein formulation suits the approach because objectives are numerous, conflicting and measurable on plate-based instruments. The



clearest closed-loop result is the three-objective antibody campaign discussed above, which reached an optimized formulation in 33 experiments while respecting osmolality and pH constraints and quantifying each excipient's contribution to each property ¹³. On the liquid side, a classification algorithm coupled to multi-objective optimization on a semi-automated robotic platform identified nine recipes meeting customer-defined criteria for a formulated liquid product within 15 working days, on a problem where no adequate physical model existed ²⁰. Long-acting injectables have been approached predictively, with a gradient boosting model trained on 181 release profiles comprising 3,783 measurements across 43 drug-polymer combinations, then validated prospectively on two newly synthesised formulations ³¹. That prospective step, rare in formulation machine learning, is what separates a model from an assertion.

6. The binding constraint is measurement, not computation

The preceding sections describe an inversion. Algorithms are adequate and improving, hardware is expensive but tractable, and the property that must be optimized frequently cannot be measured before the next decision is due.

6.1 A timescale mismatch

A campaign of thirty to a hundred experiments is efficient only if each cycle completes in hours. Particle size and zeta potential return in minutes, a dissolution profile in hours, a calorimetry trace in one to two hours. Accelerated stability at 40 degrees and 75% relative humidity takes one to six months, real-time stability one to three years, and an animal pharmacokinetic study weeks plus ethical approval no algorithm can request. Figure 4 positions the common assays against both time to result and closeness to the regulatory endpoint, and the striking feature is how empty the useful quadrant is.

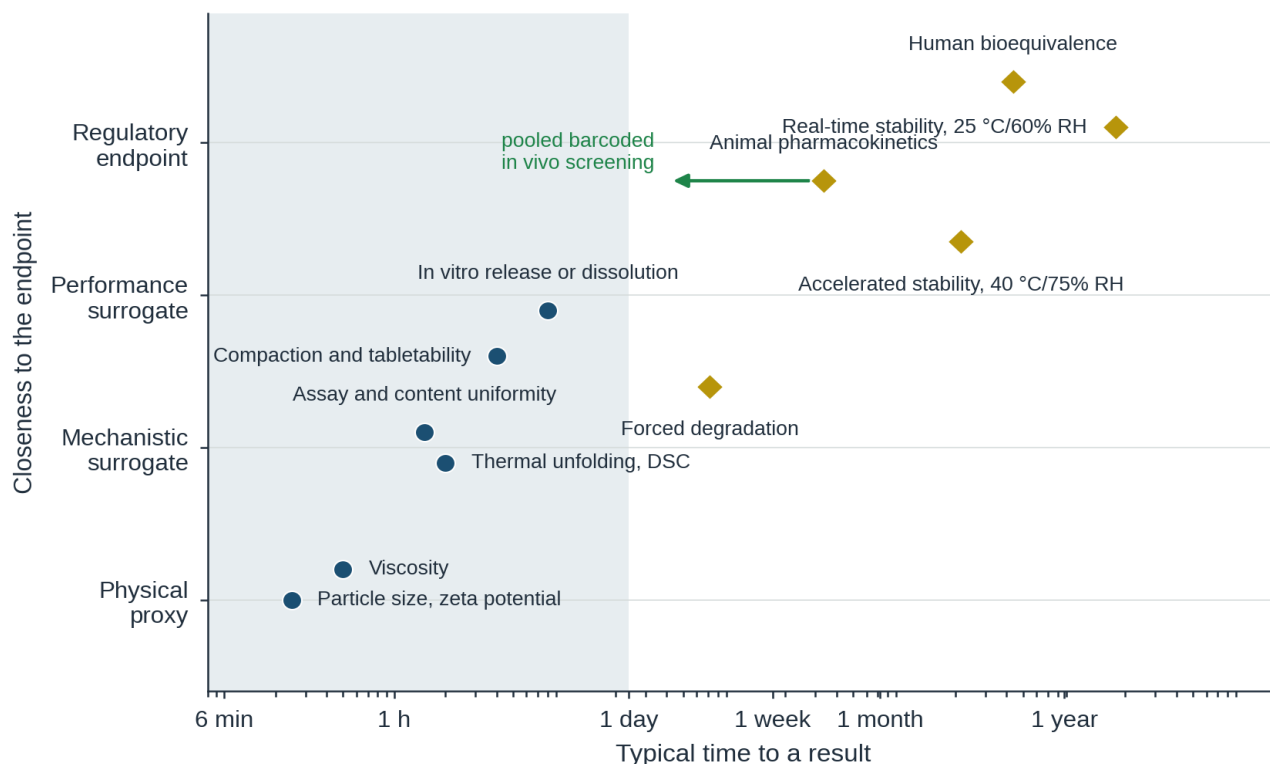


Figure 4. The measurement gap that constrains closed-loop formulation. Common pharmaceutical assays are positioned by time to result (horizontal, logarithmic, from minutes to years) and by how directly the measurement bears on the regulatory and clinical endpoint (vertical, from an indirect physical proxy to the endpoint itself). The shaded region marks assays fast enough to sit inside an optimization cycle of thirty to a hundred experiments. Assays that are both fast and endpoint-relevant are scarce, so a loop running at algorithmic speed necessarily optimizes a surrogate drawn from the lower left. The arrow marks the direction in which pooled in vivo barcoding methods have moved animal readouts, which is the pattern the rest of pharmaceuticals needs to reproduce for dissolution, physical stability and manufacturability²⁹.

The consequence is structural. Any loop that runs at the speed of the algorithm must optimize a fast surrogate, and every fast surrogate has a documented failure mode. In vitro dissolution correlates with absorption only under conditions that must themselves be established. Thermal unfolding temperature poorly predicts long-term aggregation for some antibody formats. Initial amorphous character does not determine recrystallisation risk. The loop does not know this and cannot be told it by an acquisition function.

6.2 Optimizing the surrogate

An optimizer given a proxy will exploit the proxy, including the regions where it stops tracking the endpoint. In formulation this takes a specific form: sampling concentrates near the surrogate's apparent optimum, which is precisely where a divergence would be least likely to be detected, because no diverse confirmatory sample is ever taken there. A conventional development programme is protected against this by its own inefficiency, since a formulator tests a spread of compositions out of judgement and habit and so samples the space more broadly than an optimizer would.



The failure is also hard to detect after the fact. A campaign that converges on a composition with excellent *in vitro* release and poor bioavailability produces no anomaly at the time, because every measurement the loop took was correct. The error lives in the relationship between the surrogate and the endpoint, not in any individual result, and it surfaces months later when the confirmatory study runs. Conventional development catches this earlier only because it moves slowly enough to interleave endpoint measurements with formulation work.

Three mitigations exist and all are under-used. Multi-fidelity optimization lets cheap measurements guide the search while expensive endpoint measurements correct it sparingly, though its benefit depends on how cheap and informative the low-fidelity source is ²⁴. Explicit prospective validation of a small, diverse sample of the final candidate set against the true endpoint is demonstrated in the long-acting injectable work ³¹. The third is to rebuild the assay itself, which is what barcoding did for *in vivo* screening ²⁹. That is the pattern the rest of pharmaceuticals should copy.

Physical and chemical stability is the property most resistant to closed-loop treatment and the one carrying the greatest regulatory weight. No algorithmic sophistication compresses a two-year storage study. Three responses are available: build predictive stability models with quantified uncertainty and use them as an explicit low-fidelity level; optimize a formulation's robustness

margin rather than its nominal performance, so the selected candidate has room to degrade; or accept that closed-loop methods select candidates for conventional stability assessment rather than replacing it. The third is honest, sufficient, and describes what current platforms actually do. It should be stated rather than left implicit.

6.3 The data layer

A loop is a data-generating machine, and its long-term value lies less in the campaign it completes than in the record it leaves. Formulation datasets assembled from the literature are small and biased, since successful formulations are reported and failures are not: assembled sets rarely exceed a thousand instances and are often below two hundred, so models operate permanently in the low-data regime ²⁷. A loop that discards its failures throws away the half of the record most useful to the next campaign, because the infeasible region is what the surrogate needs to learn. Pharmaceuticals has an unusually strong reason to fix this, since its records are already required to be complete, attributable and enduring under data integrity expectations ³². A regulatory dossier is a highly structured document that happens to be stored as prose.

7. Regulatory position

A formulation that cannot be filed has no value. The current position is more accommodating than is often assumed, with one substantial gap, summarised in Table 3.



Table 3. Regulatory and quality instruments bearing on closed-loop formulation development, and what each requires of an autonomous workflow.

Instrument	Status	What it requires of a closed-loop workflow	Ref.
ICH Q8(R2) Pharmaceutical Development	Step 4 guideline, current text 2009	Design space is defined by demonstrated assurance of quality, not by the experimental method used, so an adaptively derived space is admissible if the evidence is presented	³
FDA draft guidance on artificial intelligence for regulatory decision-making	Draft, January 2025	Model credibility established against a defined context of use, with evidence scaled to risk	³³
EMA reflection paper on artificial intelligence in the medicinal product lifecycle	Adopted September 2024	Risk assessed on two axes, patient risk and regulatory impact; frozen models distinguished from adaptive ones	³⁴
Data integrity expectations	Guidance 2018	Records attributable, legible, contemporaneous, original and accurate, with an audit trail permitting reconstruction of the sequence of events	³²
Analysis of machine learning in GMP environments	2025	Favours locked models validated at a fixed state with a predetermined change control protocol, and separates human-in-the-loop from human-on-the-loop accountability	³⁵

Nothing in ICH Q8(R2) requires that a design space be derived from a factorial design; the guideline defines it in terms of demonstrated assurance of quality, which is a statement about evidence rather than about method ³. A design space established by adaptive experimentation with a probabilistic model is therefore admissible in principle. An emerging perspective in the formulation literature reaches the same conclusion, arguing that intelligent optimization should sit alongside rather than replace classical design of experiments, and that acceptance depends on whether a model is credible, transparent and demonstrably fit for purpose ⁴.

Both major agencies have now stated positions. The United States Food and Drug Administration issued draft guidance in January 2025 proposing a risk-based credibility assessment framework in which the evidence required of a model scales with its context of use ³³. The European Medicines Agency adopted a reflection paper in September 2024 structured around patient risk and regulatory impact, drawing an explicit

distinction between a frozen model whose parameters are fixed and an adaptive model that continues to learn ³⁴. Neither objects to using a model to design a formulation. Both want to know what it was used for, how its credibility was established for that use, and what happens when it changes.

The unresolved problem is the autonomy rather than the model. Computerised systems in a regulated environment must produce secure, time-stamped audit trails permitting reconstruction of the sequence of events that created or modified a record ³². None of that was written for a system that selects its own next experiment. An audit trail recording what was executed does not record why, and in a closed loop the why is a posterior distribution and an acquisition score at a moment in time. Reconstructing the decision requires the model state, the training data as it stood, the random seed and the software version, none of which conventional laboratory record-keeping captures. Regulatory analysis of machine learning in good



manufacturing practice environments favours locked models validated at a fixed state with a predetermined change control protocol, and frames accountability through the distinction between human-in-the-loop operation, where a person approves each action, and human-on-the-loop operation, where the system acts under supervision³⁵. That maps onto the autonomy levels of Section 2 and suggests where the boundary will settle: level 3 autonomy with a locked decision policy and a complete decision log is defensible now, whereas a system revising its own objective is not. The practical implication is immediate. A campaign intended to support a filing should log, from the first cycle, the software version and configuration, the training set as it stood at each decision, the fitted model or its serialised state, the acquisition values that ranked the candidates, the random seed, the condition actually executed as distinct from the condition proposed, and the raw analytical output rather than only the derived value. None of that is expensive to capture at the time and none of it can be reconstructed afterwards. A laboratory that treats the decision log as an afterthought will find that its most efficient campaign is also the one it cannot defend.

8. Priorities for the field

Five programmes would change the state of this field more than further algorithmic refinement.

Report sample efficiency against a stated baseline. Every closed-loop formulation paper should give the size of the design space, the number of experiments consumed, and either a

head-to-head comparison with a conventional design or a random-sampling control on the same space. The metallophotocatalyst campaign is the model to follow¹⁸. Acceleration factors quoted without a denominator should be read as marketing.

Build fast surrogate assays with characterised relationships to regulatory endpoints. This is the highest-value experimental programme available. Nanomedicine has shown what it looks like when a slow endpoint is re-engineered for throughput²⁹; the equivalent for dissolution, physical stability and manufacturability does not exist.

Solve solid handling. Reliable low-mass dispensing of cohesive powders is a prerequisite for autonomous work on the dosage forms that constitute most of the market, and it has been named as an open bottleneck¹⁵. It is an engineering problem rather than a scientific one, which is precisely why it is neglected.

Publish failures in machine-readable form. A closed-loop campaign generates a complete record of what did not work, and that record is more informative than the optimum. The infeasible region is what makes the next campaign efficient²⁷.

Design the audit trail before the campaign, capturing model state and decision rationale alongside execution^{32,35}. Building this into orchestration software from the start is cheap; adding it afterwards is not.



9. Conclusion

Closed-loop formulation development is real, it works, and its published results in pharmaceuticals are now specific enough to assess rather than merely to anticipate. Bayesian optimization reliably reaches good formulations with fewer experiments than a fixed design, particularly where objectives conflict and constraints are hard, and robotic execution has carried a full tableting workflow from raw material to specification-compliant product within a working day. What has not been solved, and what no computational progress will solve on its own, is that the properties defining a medicine are measured on timescales the loop cannot accommodate, so every autonomous campaign optimizes something other than the property that matters. The next decade will be decided less by better algorithms than by better assays, by the discipline of reporting efficiency against real baselines, and by whether pharmaceuticals finally builds the machine-readable record of its own experiments that the loop needs in order to learn.

Declarations

Conflicts of interest: None declared.

Funding: No specific funding was received for this work.

References

1. Bannigan P, Hickman RJ, Aspuru-Guzik A, Allen C. The dawn of a new pharmaceutical epoch: can AI and robotics reshape drug formulation? *Adv Healthc Mater.* 2024;13(29):2401312. doi:10.1002/adhm.202401312
2. Wouters OJ, McKee M, Luyten J. Estimated research and development investment needed to bring a new medicine to market, 2009-2018. *JAMA.* 2020;323(9):844-53. doi:10.1001/jama.2020.1166
3. International Council for Harmonisation of Technical Requirements for Pharmaceuticals for Human Use. ICH harmonised tripartite guideline Q8(R2): pharmaceutical development. Geneva: ICH; 2009. Available from: https://database.ich.org/sites/default/files/Q8_R2_Guideline.pdf
4. Turkdogan A, Alloush T, Demiralp B. Experimental design in pharmaceutical formulation development: achievements, limitations and the transition toward intelligent optimization. *Sci Pharm.* 2026;94(2):38. doi:10.3390/scipharm94020038
5. Jones DR, Schonlau M, Welch WJ. Efficient global optimization of expensive black-box functions. *J Glob Optim.* 1998;13(4):455-92. doi:10.1023/A:1008306431147
6. Shahriari B, Swersky K, Wang Z, Adams RP, de Freitas N. Taking the human out of the loop: a review of Bayesian optimization. *Proc IEEE.* 2016;104(1):148-75. doi:10.1109/JPROC.2015.2494218
7. Shields BJ, Stevens J, Li J, et al. Bayesian reaction optimization as a tool for chemical



- synthesis. *Nature*. 2021;590(7844):89-96. doi:10.1038/s41586-021-03213-y
8. Sano S, Kadowaki T, Tsuda K, Kimura S. Application of Bayesian optimization for pharmaceutical product development. *J Pharm Innov*. 2020;15(3):333-43. doi:10.1007/s12247-019-09382-8
9. Tom G, Schmid SP, Baird SG, et al. Self-driving laboratories for chemistry and materials science. *Chem Rev*. 2024;124(16):9633-732. doi:10.1021/acs.chemrev.4c00055
10. Burger B, Maffettone PM, Gusev VV, et al. A mobile robotic chemist. *Nature*. 2020;583(7815):237-41. doi:10.1038/s41586-020-2442-2
11. Abolhasani M, Kumacheva E. The rise of self-driving labs in chemical and materials sciences. *Nat Synth*. 2023;2:483-92. doi:10.1038/s44160-022-00231-0
12. Tobias AV, Wahab A. Autonomous 'self-driving' laboratories: a review of technology and policy implications. *R Soc Open Sci*. 2025;12(7):250646. doi:10.1098/rsos.250646
13. Waibel I, Schneider TN, Fischer FJ, et al. Bayesian optimization for efficient multiobjective formulation development of biologics. *Mol Pharm*. 2025;22(11):6636-45. doi:10.1021/acs.molpharmaceut.5c00591
14. Gardner JR, Kusner MJ, Xu Z, Weinberger KQ, Cunningham JP. Bayesian optimization with inequality constraints. In: *Proceedings of the 31st international conference on machine learning*. PMLR 32. 2014. p. 937-45. Available from: <https://proceedings.mlr.press/v32/gardner14.html>
15. Seifrid M, Pollice R, Aguilar-Granda A, et al. Autonomous chemical experiments: challenges and perspectives on establishing a self-driving lab. *Acc Chem Res*. 2022;55(17):2454-66. doi:10.1021/acs.accounts.2c00220
16. Hase F, Aldeghi M, Hickman RJ, Roch LM, Aspuru-Guzik A. Gryffin: an algorithm for Bayesian optimization of categorical variables informed by expert knowledge. *Appl Phys Rev*. 2021;8(3):031406. doi:10.1063/5.0048164
17. Hickman RJ, Sim M, Pablo-Garcia S, et al. Atlas: a brain for self-driving laboratories. *Digit Discov*. 2025;4(4):1006-29. doi:10.1039/D4DD00115J
18. Li X, Che Y, Chen L, et al. Sequential closed-loop Bayesian optimization as a guide for organic molecular metallophotocatalyst formulation discovery. *Nat Chem*. 2024;16(8):1286-94. doi:10.1038/s41557-024-01546-5
19. Ros H, Abdalla Y, Cook MT, Shorthouse D. Efficient discovery of new medicine formulations using a semi-self-driven robotic formulator. *Digit Discov*. 2025;4:2263-72. doi:10.1039/D5DD00171D
20. Cao L, Russo D, Felton K, et al. Optimization of formulations using robotic experiments driven by machine learning



- DoE. Cell Rep Phys Sci. 2021;2(1):100295. doi:10.1016/j.xcrp.2020.100295
- 2024;635(8040):890-7. doi:10.1038/s41586-024-08173-7
21. Xu Y, Ma S, Cui H, et al. AGILE platform: a deep learning powered approach to accelerate LNP development for mRNA delivery. Nat Commun. 2024;15:6305. doi:10.1038/s41467-024-50619-z
 22. Li B, Raji IO, Gordon AGR, et al. Accelerating ionizable lipid discovery for mRNA delivery using machine learning and combinatorial chemistry. Nat Mater. 2024;23(7):1002-8. doi:10.1038/s41563-024-01867-3
 23. Volk AA, Epps RW, Yonemoto DT, et al. AlphaFlow: autonomous discovery and optimization of multi-step chemistry using a self-driven fluidic lab guided by reinforcement learning. Nat Commun. 2023;14(1):1403. doi:10.1038/s41467-023-37139-y
 24. Sabanza-Gil V, Barbano R, Pacheco Gutierrez D, et al. Best practices for multi-fidelity Bayesian optimization in materials and molecular research. Nat Comput Sci. 2025;5(7):572-81. doi:10.1038/s43588-025-00822-9
 25. Abbas F, Salehian M, Hou P, et al. Accelerated drug development using a digital formulator and a self-driving tableting data factory. Nat Commun. 2026;17:4739. doi:10.1038/s41467-026-71204-6
 26. Dai T, Vijayakrishnan S, Szczypinski FT, et al. Autonomous mobile robots for exploratory synthetic chemistry. Nature. 2024;635(8040):890-7. doi:10.1038/s41586-024-08173-7
 27. Hickman RJ, Bannigan P, Bao Z, Aspuru-Guzik A, Allen C. Self-driving laboratories: a paradigm shift in nanomedicine development. Matter. 2023;6(4):1071-81. doi:10.1016/j.matt.2023.02.007
 28. Canty RB, Bennett JA, Brown KA, et al. Science acceleration and accessibility with self-driving labs. Nat Commun. 2025;16:3856. doi:10.1038/s41467-025-59231-1
 29. Rhym LH, Manan RS, Koller A, Stephanie G, Anderson DG. Peptide-encoding mRNA barcodes for the high-throughput in vivo screening of libraries of lipid nanoparticles for mRNA delivery. Nat Biomed Eng. 2023;7(7):901-10. doi:10.1038/s41551-023-01030-4
 30. Muniz Castro B, Elbadawi M, Ong JJ, et al. Machine learning predicts 3D printing performance of over 900 drug delivery systems. J Control Release. 2021;337:530-45. doi:10.1016/j.jconrel.2021.07.046
 31. Bannigan P, Bao Z, Hickman RJ, et al. Machine learning models to accelerate the design of polymeric long-acting injectables. Nat Commun. 2023;14:35. doi:10.1038/s41467-022-35343-w
 32. US Food and Drug Administration. Data integrity and compliance with drug CGMP: questions and answers, guidance for industry. Silver Spring (MD): FDA; 2018. Available from:



<https://www.fda.gov/media/119267/download>

33. 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 and other interested parties. Docket FDA-2024-D-4689. Silver Spring (MD): FDA; 2025. Available from: <https://www.federalregister.gov/documents/2025/01/07/2024-31542/>
34. European Medicines Agency. Reflection paper on the use of artificial intelligence in the medicinal product lifecycle. EMA/CHMP/CVMP/83833/2023. Amsterdam: EMA; 2024. Available from: https://www.ema.europa.eu/en/documents/scientific-guideline/reflection-paper-use-artificial-intelligence-ai-medicinal-product-lifecycle_en.pdf
35. Niazi SK. Regulatory perspectives for AI/ML implementation in pharmaceutical GMP environments. *Pharmaceuticals* (Basel). 2025;18(6):901. doi:10.3390/ph18060901