



Quantitative Bridge or Black Box? A Practice-Based Framework for Statistical Programming Contributions to Population Pharmacokinetic/Pharmacodynamic Modeling in Regulatory Drug Development

Radhika Aitha

Independent Researcher, USA

ORCID ID: 0009-0006-9213-0700

KEYWORDS:

Population pharmacokinetics; Model-informed drug development; Statistical programming; Noncompartmental analysis; CDISC ADaM; Regulatory submissions.

Corresponding Author:

Radhika Aitha

DOI: [10.55677/IJMSPR/2026-3050-I908](https://doi.org/10.55677/IJMSPR/2026-3050-I908)

Published: September 28, 2026

License:

This is an open access article under the CC BY 4.0 license:
<https://creativecommons.org/licenses/by/4.0/>

ABSTRACT

Population pharmacokinetic/pharmacodynamic (PK/PD) modeling now sits at the center of model-informed drug development (MIDD), shaping dose selection, labeling language, and regulatory submissions across therapeutic areas. Less attention goes to the statistical and clinical programmers who build the analysis-ready datasets. These models depend on a role distinct from the pharmacometrician who fits the model itself. This article proposes a four-domain framework for describing that contribution: noncompartmental analysis (NCA) dataset construction and quality control; population PK dataset and diagnostic-output support; exposure-response and dose-selection analysis support; and specialized regulatory analyses spanning drug-drug interaction, bioequivalence, and pediatric extrapolation studies. Drawing on FDA and CDISC guidance, implementation standards, and pharmacometrics practice literature, the synthesis characterizes the technical demands, regulatory touchpoints, and quality-control conventions specific to each domain. Programmer contribution turns out to escalate unevenly rather than proportionally across domains: dataset-construction choices especially around below-limit-of-quantification handling and covariate dataset structure carry direct downstream weight on model validity and labeling claims, and this weight increases sharply once analyses move into specialized regulatory territory. With FDA's PDUFA VII commitments extending MIDD infrastructure through fiscal years 2023 to 2027, the framework gives a vocabulary for a programming specialization that current workforce literature treats as undersupplied relative to demand, with direct bearing on training pathways and staffing models for regulatory-facing quantitative pharmacology teams.

Cite the Article: Aitha, R. (2026). Quantitative Bridge or Black Box? A Practice-Based Framework for Statistical Programming Contributions to Population Pharmacokinetic/Pharmacodynamic Modeling in Regulatory Drug Development. *International Journal of Medical Science and Pharmaceutical Research*, 3(9), 510-518. <https://doi.org/10.55677/IJMSPR/2026-3050-I907>

INTRODUCTION

Model-informed drug development has stopped being a supplementary analytical add-on and become a formal regulatory expectation. FDA's PDUFA VII commitment letter, released in August 2021 and covering fiscal years 2023 through 2027, reinforces that shift by extending MIDD infrastructure and review capacity into the current review cycle.¹³ Population PK/PD modeling sits at the center of the resulting work: FDA finalized its current Population Pharmacokinetics guidance for industry in February 2022, formally establishing the agency's expectations for population PK data and modeling submitted as part of new drug and biologics applications.^{1,2} It now routinely justifies dose selection, supports labeling language, and answers regulatory questions that would otherwise demand additional dedicated trials.¹

None of that modeling happens without a dataset infrastructure someone else builds, validates, and maintains typically a statistical or clinical programmer working from Clinical Data Interchange Standards Consortium (CDISC) Study Data Tabulation Model (SDTM) and Analysis Data Model (ADaM) standards.¹² That division of labor matters operationally, but the literature rarely describes it with any precision. Workforce commentary on the MIDD skills gap focuses almost entirely on modelers and quantitative pharmacologists,¹⁴ and the programming contribution goes largely untheorized despite its direct bearing on whether a model can be trusted at all.

This gap has practical consequences beyond academic tidiness. When a sponsor organization plans staffing for a PK/PD-heavy submission, the planning conversation tends to center on how many pharmacometricians the program needs, with programming support treated as an undifferentiated extension of general clinical programming capacity. That framing obscures real differences in what PK/PD programming actually requires from one task to the next, and it leaves training investment concentrated on modeling skills while the dataset-construction skills the models depend on receive comparatively little dedicated attention.

The distinction between programmer and pharmacometrician is not a matter of formal job title alone; it reflects two genuinely different professional orientations toward the same underlying data. The pharmacometrician's primary responsibility is specifying and interpreting a model choosing a structural form, deciding which covariates merit testing, and interpreting what a fitted parameter implies about drug behavior in the population under study. The programmer's primary responsibility is ensuring that the dataset the model is fit to accurately and completely represents the underlying clinical data, correctly encodes the dosing and sampling history, and is structured in a way the modeling software can consume without silent misinterpretation. These are complementary skills, not interchangeable ones, and a development program that treats them as a single undifferentiated "PK/PD analysis" function risks understaffing whichever half of that function is less visible in a given organization's reporting structure.

This distinction matters more, not less, as MIDD approaches expand beyond their traditional home in clinical pharmacology groups into broader use across a development program. As exposure-based modeling is applied to an increasing range of regulatory questions from traditional dose selection to newer applications in dose optimization for oncology products and model-based bridging across formulations or populations the volume of dataset-construction work required to support this modeling grows correspondingly, and the professional distinction between programmer and pharmacometrician becomes more, not less, operationally significant to get right in workforce planning.

This article works to close that gap. It proposes a four-domain framework for statistical programming contributions to PK/PD analysis: noncompartmental analysis dataset construction, population PK dataset and diagnostic support, exposure-response and dose-selection analysis support, and specialized regulatory analyses. Each domain carries its own technical conventions, quality-control demands, and regulatory exposure. The framework's central claim is that these demands do not scale uniformly with study complexity they escalate unevenly, and knowing where the sharpest increases occur has practical value for how programming teams are trained and staffed.

The remainder of this article proceeds as follows. The Materials and Methods section describes the literature-synthesis approach used to construct the framework and the criteria applied to source selection. The Results section presents each of the four domains in turn, characterizing the technical conventions, regulatory touchpoints, and dataset-construction demands specific to each, and closes with a comparative matrix and escalation diagram summarizing how these demands relate to one another. The Discussion section addresses the practical implications of the uneven-escalation finding for MIDD workforce planning, training design, and staffing structure within CROs and sponsor organizations, before closing with the limitations of a literature-synthesis approach to characterizing professional practice.

MATERIALS AND METHODS

This article synthesizes publicly available regulatory guidance documents, CDISC standards documentation, and peer-reviewed pharmacometrics practice literature describing statistical programming contributions to PK/PD analysis. Sources were located through targeted search of FDA guidance repositories, the CDISC standards library, and peer-reviewed pharmacometrics journals, including CPT: Pharmacometrics & Systems Pharmacology and Clinical Pharmacology & Therapeutics, with priority given to current final guidance documents and recently published practice literature over superseded drafts. Where a guidance document had been through multiple revisions as with the population pharmacokinetics guidance, finalized in 2022 after a 2019 draft, or the cytochrome P450 drug-interaction guidance, finalized in 2020 after a 2017 draft the current final version was used as the authoritative source, with the draft-to-final revision history noted where it clarified how the agency's expectations had shifted.

Twenty-one sources met relevance and currency criteria and were organized against the four-domain framework described above: noncompartmental analysis and below-limit-of-quantification handling; population pharmacokinetic modeling, CDISC ADaM implementation, and model-evaluation diagnostics; exposure-response analysis and dose-selection guidance; and drug-drug interaction, bioequivalence, and pediatric extrapolation methodology, together with workforce and regulatory-commitment literature framing the discussion. Sources were weighted toward primary regulatory guidance and peer-reviewed methodology papers over secondary commentary, on the reasoning that a practice framework intended to describe programmer-facing regulatory demands should be anchored as directly as possible in the documents that actually define those demands, rather than in

interpretive summaries of them.

Each domain was assessed for its characteristic technical conventions, the regulatory touchpoint at which programmer-constructed outputs are directly scrutinized, and the degree to which its dataset conventions transfer to other domains. Technical conventions were identified from the specific procedural language used in guidance documents and implementation guides for example, the explicit BLQ-handling rules documented in the noncompartmental analysis literature, or the specific covariate-modeling extensions documented in the ADaM popPK Implementation Guide. Regulatory touchpoints were identified by tracing where each domain's output is explicitly referenced in a regulatory decision context, whether a labeling claim, a dose-justification argument, or a specific study-type submission requirement. Transferability was assessed qualitatively, by comparing whether the dataset conventions, source data types, and quality-control logic documented for one domain appeared, on inspection of the underlying guidance and methodology literature, to carry over into an adjacent domain or to require substantively new conventions.

This comparative structure, rather than a statistical test, forms the basis of the framework presented in the Results section. The four dimensions in Table 1 primary deliverable, regulatory touchpoint, and relative technical and regulatory intensity were derived from this qualitative comparison rather than from a quantitative scoring instrument, and the ordinal intensity ratings should be read as a structured summary of the literature comparison rather than as an independently validated measurement scale.

This is a literature-synthesis and practice-framework article. It does not involve original patient data, human subjects research, or animal research, so an ethics committee approval statement does not apply here. No primary or client-derived dataset was analyzed every figure cited traces back to the reference literature at its point of use, and no numerical value presented in this article was estimated, extrapolated, or invented beyond what its cited source directly states or what is explicitly labeled as an illustrative, author-constructed synthesis for expository purposes.

RESULTS

Domain 1 Noncompartmental Analysis Dataset Construction

Noncompartmental analysis (NCA) is still the first analytical layer applied to concentration-time data in early-phase PK studies, producing the derived parameters maximum observed concentration, time to maximum concentration, area under the concentration-time curve, terminal half-life that feed downstream population modeling.⁵ Unlike compartmental approaches, NCA does not assume a specific structural model of drug disposition; it estimates exposure directly from the observed concentration-time profile using the trapezoidal rule and related numerical integration methods. That apparent simplicity is precisely what makes the underlying dataset construction consequential: because there is no model to compensate for a poorly constructed input, every derived parameter is only as reliable as the concentration-time dataset feeding it.

Below-limit-of-quantification (BLQ) handling turned out to be one domain-specific convention with outsized downstream consequence. BLQ values before the first measurable concentration are conventionally set to zero; BLQs after the last measurable concentration are set to missing; single BLQ values embedded between two measurable concentrations are also set to missing, as are consecutive embedded BLQs; and profiles that are entirely BLQ get excluded from analysis altogether.⁵ Each of these rules encodes a specific assumption about drug disposition for example, that a BLQ value preceding the first measurable concentration reflects a genuine pre-dose or absorption-lag state rather than an assay failure and applying the wrong rule to an ambiguous profile silently distorts the resulting AUC or half-life estimate rather than producing an obvious error the programmer or modeler would catch downstream. More sophisticated alternatives exist, including maximum-likelihood-based imputation methods and kernel-density-based estimation approaches for BLQ-affected AUC calculation, but these methods are used inconsistently across organizations and require the programmer to understand the statistical rationale behind whichever convention a given sponsor has adopted, not merely to execute a fixed rule mechanically.

These are decisions the programmer executes at the dataset level, not choices made inside the model, which means inconsistent application threatens the validity of parameters that later show up in regulatory submissions, often without the modeler ever seeing where the inconsistency originated. A modeler working from a derived NCA parameter dataset generally has no visibility into which BLQ convention was applied to which subject unless the programmer has documented that decision explicitly in the dataset metadata which places a documentation burden on Domain 1 work that is easy to underestimate when resourcing plans treat NCA support as a comparatively lightweight, junior-appropriate task.

The CDISC ADaM standard provides the structural scaffolding within which these decisions are implemented. In an SDTM database, pharmacokinetic data is stored as one record per subject per time point (the PC domain) or per derived parameter (the PP domain); analysis datasets (ADPC, ADPP) compiled from these SDTM domains together with subject-level ADaM data must correctly link dosing and sampling records, typically through a formal relational table, before any derived parameter can be trusted.^{11,12} This linkage step is deceptively simple to describe and easy to get wrong in practice: dosing and sampling times are frequently captured by different clinical staff using different systems, and a programmer reconciling the two must resolve discrepancies a sample drawn slightly before or after its nominal protocol-specified time, a dose recorded to the nearest hour when the sampling schedule requires minute-level precision using judgment that draws on both the protocol's intent and the pharmacokinetic plausibility of the resulting profile.

Programmer judgment at this stage, particularly around anomalous or embedded BLQ values within an otherwise measurable profile, functions as a quality gate that later analytical stages assume has already been passed correctly. Because Domain 1 output feeds directly into Domain 2 population modeling in most development programs, an unresolved ambiguity at this stage does not stay contained; it propagates into covariate datasets, model diagnostics, and ultimately into any labeling claim the population model is used to support, which is why this domain's apparent technical simplicity does not translate into low regulatory consequence.

Domain 2 Population PK Dataset and Diagnostic-Output Support

Population PK modeling, formalized in FDA's 2022 final guidance finalizing the agency's 2019 draft, runs on nonlinear mixed-effects structures that separate fixed effects population-level covariate relationships such as body weight or renal function from random effects, which capture between- and within-subject variability not explained by measured covariates.^{1,18} The random-effects component is typically assumed to follow a multivariate normal distribution, an assumption that, when misspecified, can bias both the population parameter estimates and any subsequent inference about individual variability, a statistical subtlety that has direct dataset implications, since the covariate structure a programmer builds determines which sources of variability the model can even attempt to explain versus which get absorbed into the random-effects term as unexplained noise.

CDISC's ADaM popPK Implementation Guide, a comparatively recent standard relative to the broader ADaM framework, defines the Basic Data Structure extensions this modeling needs, and it falls to the programmer to build covariate datasets and link PK concentration and parameter domains through structured relational tables.^{12,13} Constructing a covariate dataset for population PK analysis is materially more demanding than constructing a standard efficacy or safety analysis dataset, because covariates must be time-varying where clinically appropriate renal function measured at multiple visits, concomitant medication use that starts and stops during the study rather than captured as a single baseline value. A programmer who defaults to baseline-only covariate handling out of habit from safety-dataset conventions can silently understate a clinically meaningful covariate relationship the modeler was specifically trying to characterize.

Model evaluation diagnostics goodness-of-fit plots comparing observed to population- and individual-predicted concentrations, conditional weighted residual plots, and prediction-corrected visual predictive checks that account for heterogeneous dosing and sampling designs all depend on programmer-built output datasets structured to a modeler's specification.¹⁷ Producing a valid prediction-corrected visual predictive check in particular requires the programmer to understand not just the mechanics of percentile calculation but why a standard, uncorrected visual predictive check becomes misleading when a study design includes adaptive dosing or substantially varied sampling schedules across subjects; the correction exists specifically to prevent an artifact of study design from masquerading as a genuine model misfit.

In practice, this makes the programmer a direct determinant of whether model diagnostics can be produced at all, not merely a support function running behind the scenes. A dataset missing a covariate the modeler needs for a planned diagnostic, or structured with an incompatible time variable, can stall a model-qualification timeline regardless of how sound the underlying model specification is. Internal guidance documents developed by individual sponsor organizations to standardize this handoff illustrate how much operational weight rests on getting the dataset specification right the first time, rather than treating it as a formality ahead of the modeler's "real" work.¹⁶ Organizations that have documented their own internal population PK dataset-preparation practices report that establishing this specification collaboratively, with the programmer involved before the modeling plan is finalized rather than receiving it as a fixed handoff, meaningfully reduces the number of dataset-revision cycles required once modeling begins an operational detail that rarely appears in published methodology but shapes how efficiently Domain 2 work actually proceeds inside a program.

Domain 3 Exposure-Response and Dose-Selection Analysis Support

FDA's guidance on exposure-response relationships pushes sponsors to start dose- and exposure-response evaluation early in development, updating continuously as data accumulate, precisely because poor early dose selection is difficult and costly to correct later in a program.¹⁹ The guidance draws a deliberate distinction between exploratory exposure-response analyses conducted during early development, which can tolerate looser dataset specifications because their purpose is to inform the design of later, more definitive studies, and confirmatory analyses conducted closer to submission, which require the same dataset rigor expected of any pivotal efficacy or safety analysis. A programmer moving between these two modes within the same program must recognize which standard currently applies, since applying exploratory-stage dataset shortcuts to a confirmatory-stage analysis risks understating the analysis's regulatory defensibility.

Programmer work in this domain centers on building exposure metrics steady-state area under the curve, maximum concentration, and trough concentration, depending on which best correlates with the endpoint of interest linked to efficacy and safety endpoints, and FDA guidance on alternative dosing regimens draws directly on this work for labeling-relevant dose justification.³ Constructing these linked datasets requires the programmer to merge two data streams that are rarely collected on the same schedule: PK sampling, which follows a schedule designed to characterize the concentration-time profile, and endpoint assessment, which follows a schedule designed around clinical or laboratory assessment logistics. Deriving a defensible exposure metric at the time of endpoint assessment rather than simply pairing the nearest available PK sample regardless of temporal distance often requires model-based

interpolation informed by the Domain 2 population PK model, meaning Domain 3 work frequently depends on Domain 2 output being both available and structured compatibly, a dependency that resourcing plans treating each domain as an independent task can easily overlook.

Poor dose selection is documented as one of several recurring drivers of Phase III trial failure, alongside issues in basic science, clinical study design, and data collection and analysis, which raises the practical stakes of getting these exposure-metric datasets right well before a confirmatory trial begins.²¹ One recurring pattern in the trial-failure literature is that a dose carried forward into Phase III was selected based on a Phase II exposure-response analysis built on a comparatively thin or poorly characterized exposure dataset not because the underlying pharmacology was misunderstood, but because the dataset supporting the dose-selection analysis did not adequately capture the exposure variability present in the studied population. A dose-selection dataset built on an incompletely characterized exposure metric does not simply produce a weaker analysis; it can propagate an avoidable dosing error into a program's most expensive and most difficult-to-repeat trials, and by the time the consequence becomes visible typically an unexpectedly narrow therapeutic margin or an efficacy shortfall in the confirmatory trial the dataset-construction decision that contributed to it is several years and several million dollars in the past.

Domain 4 Specialized Regulatory Analyses

Three sub-areas carry programming demands that go beyond general PK support, and each has developed its own dataset conventions largely independently of the others. Drug-drug interaction (DDI) analyses, governed by FDA's 2020 final guidance on cytochrome P450 enzyme- and transporter-mediated drug interactions, require datasets built to support a systematic, risk-based DDI evaluation framework that classifies a drug's interaction potential according to its metabolic and transporter profile before any dedicated interaction study is even designed.⁷⁻⁸ A programmer supporting DDI analysis must construct datasets that correctly pair perpetrator and victim drug administration records, often across a crossover or fixed-sequence design where the same subject serves as their own control, and must track washout periods precisely enough that carryover effects from a prior treatment period can be identified and, where necessary, statistically addressed rather than assumed away.

Bioequivalence analyses typically run on two-sequence, two-period crossover designs or four-period replicate designs for highly variable drugs and need programmer-built datasets that support log-transformed exposure comparisons inside FDA's confidence-interval-based bioequivalence framework.⁹⁻¹⁰ The log-transformation requirement is not a cosmetic statistical preference; it reflects the multiplicative, rather than additive, nature of the biological variability typically observed in exposure measures, and a programmer who applies an additive-scale confidence interval to what should be a log-scale comparison will generate a bioequivalence conclusion that does not match what the regulatory framework actually requires. Replicate designs, used specifically for highly variable drugs whose within-subject variability would otherwise make a standard two-period design impractically large, introduce an additional dataset complexity: the programmer must track within-subject replicate measurements across four periods per subject rather than two, and must correctly attribute variability components to within-subject versus between-subject sources when supporting the reference-scaled average bioequivalence approach that this design enables.

Pediatric extrapolation analyses lean on allometric scaling of clearance and other PK parameters from adult data. FDA's own decision-tree framework treats simple weight-based allometric scaling as scientifically justified for children roughly two years of age and older, while favoring physiologically based pharmacokinetic approaches below that threshold, where allometry's interpolative power is weaker.^{4-5,6} Constructing the dataset that supports this scaling exercise requires the programmer to source and validate reference body-weight and, where relevant, organ-maturation data for the pediatric population under study, then apply the chosen scaling exponent commonly three-quarters power for clearance, consistent with allometric principles observed across mammalian species to translate an adult PK parameter distribution into a predicted pediatric distribution. Retrospective evaluations of this approach across multiple drugs have found that simple allometric scaling predicts pediatric clearance within a twofold range of observed values in most cases, but with meaningfully lower accuracy for drugs eliminated through pathways that mature at a different rate than body size during early childhood, which means the programmer's dataset must also capture which elimination pathway applies before the scaling exponent can be selected with confidence.

None of these three sub-areas uses dataset conventions interchangeable with general PK programming each demands its own specialized build, and a programmer moving between them cannot simply carry forward the conventions learned in Domains 1 through 3, or, in most cases, from one Domain 4 sub-area to another. A programmer experienced in DDI dataset construction does not automatically know how to construct a replicate-design bioequivalence dataset, and neither skill transfers directly to pediatric allometric-scaling support; each sub-area has its own governing guidance document, its own dataset architecture, and its own failure modes that only become visible through direct experience with that specific analysis type.

Table 1: Comparative technical and regulatory demands across the four PK/PD programming domains (author's synthesis of cited sources).

Domain	Primary Deliverable	Regulatory Touchpoint	Relative Intensity
1. Noncompartmental Analysis	Concentration-time datasets; derived PK parameters (C _{max} , T _{max} , AUC, t _{1/2})	Foundational feeds all downstream modeling	Low-moderate
2. Population PK Support	Covariate datasets; PC/PP-linked ADaM structures; GOF/VPC diagnostic outputs	Model qualification for regulatory submission	Moderate-high
3. Exposure-Response Support	Exposure metric datasets (steady-state AUC, C _{max}) linked to endpoints	Dose selection and labeling justification	High
4. Specialized Regulatory Analyses	DDI, bioequivalence, and pediatric extrapolation datasets	Direct evidentiary basis for specific regulatory questions	Highest

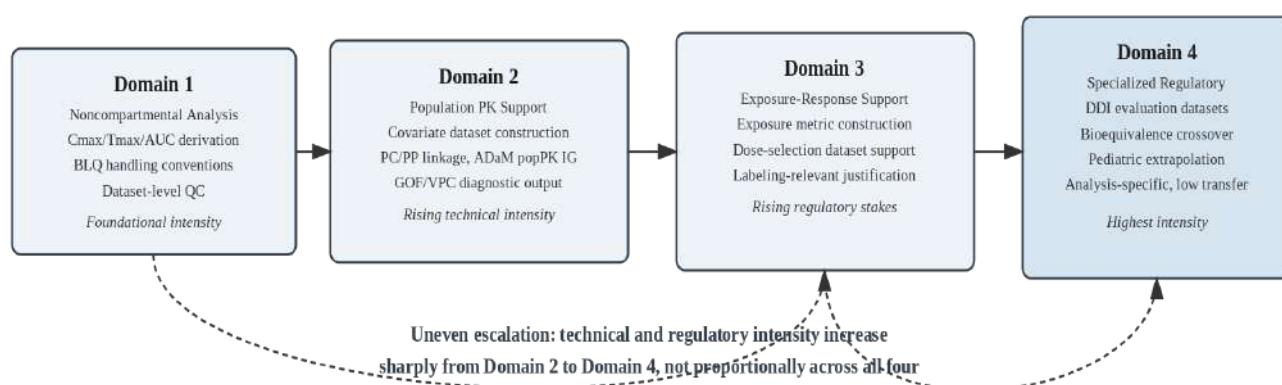


Figure basis: FDA population PK (2022) and DDI (2020) guidance; CDISC ADaM popPK IG; Jansen et al. (2018) VPC methodology.
Illustrative synthesis diagram, not a measured process-mining output.

Figure 1: Escalating technical and regulatory intensity across the four PK/PD programming domains, with the sharpest increase occurring at the transition into specialized regulatory analyses.

DISCUSSION

The four-domain framework suggests programming contribution to PK/PD analysis does not scale uniformly with study complexity. It escalates unevenly, and the sharpest jump in regulatory consequence happens at the transition from general population PK support (Domain 2) to specialized regulatory analyses (Domain 4), where dataset conventions turn highly analysis-specific and stop transferring cleanly across studies. This unevenness has a practical corollary that is easy to miss when programming capability is assessed only in terms of years of experience: a programmer with a decade of general population PK experience is not thereby prepared for pediatric extrapolation or replicate-design bioequivalence work, because the escalation described here is not primarily a matter of accumulated general skill but of exposure to a specific, comparatively narrow set of regulatory conventions that Domain 4 sub-areas do not share with each other or with the broader PK programming discipline.

The documented shortage of MIDD-trained personnel is addressed mostly through modeler- and pharmacometrician-focused training programs, including newly launched graduate curricula specifically built around quantitative systems pharmacology,^{14,20} which leaves the programming specialization this framework describes largely unaddressed in workforce planning. This asymmetry is understandable given how MIDD workforce discussions are typically framed: the visible bottleneck is usually described as too few scientists capable of building and interpreting the models themselves, since a shortage there directly constrains how many MIDD-supported regulatory interactions a sponsor or agency can conduct. What this framing tends to obscure is that a modeler's capacity is itself bounded by the availability of correctly constructed input datasets, and if the programming capacity feeding that pipeline is treated as generically available clinical programming staff rather than as a domain-specific specialization requiring its own training investment, the modeling bottleneck may simply be masking a programming bottleneck one step upstream.

PDUFA VII formally established a dedicated MIDD paired-meeting program between FDA's drug and biologics review centers and industry sponsors for fiscal years 2023 through 2027,¹⁵ and commits FDA to continued MIDD expansion through that window,¹³ so the specialized programming capacity this article maps is positioned to become a scarcer and more consequential input to drug

development timelines, not a less important one. As MIDD-supported regulatory interactions become more routine rather than exceptional, the volume of Domain 2 through Domain 4 programming work required to support them scales accordingly, and organizations that have not built a deliberate pipeline for developing this specialization as distinct from a general clinical programming pipeline may find that their capacity to participate in MIDD-supported regulatory pathways is constrained less by modeling expertise than by the availability of programmers who can reliably deliver model-ready datasets for each of the four domains described here.

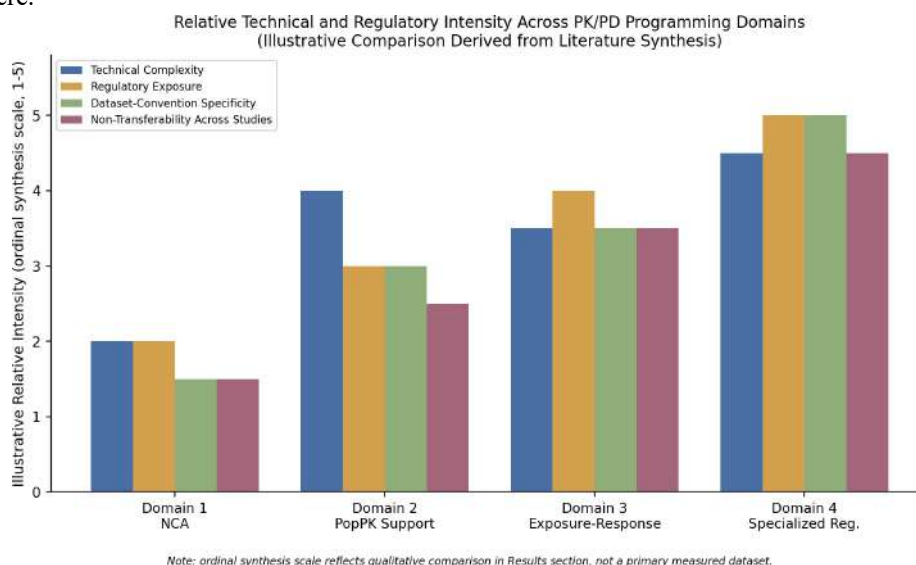


Figure 2: Relative technical and regulatory intensity across PK/PD programming domains illustrative ordinal synthesis derived from the literature comparison in Results, not a primary measured dataset.

That has real implications for how contract research organizations and sponsor organizations structure their programming teams. Treating PK/PD programming as an undifferentiated extension of general clinical programming risks under-resourcing precisely the domains specialized regulatory analyses most of all that carry the highest downstream regulatory exposure. A programmer well-versed in Domain 1 and Domain 2 conventions is not automatically prepared for a bioequivalence crossover dataset or a pediatric allometric-scaling analysis; treating that transition as a matter of general seniority rather than specific domain exposure understates what the work actually requires.

The practical staffing question this raises is not simply whether to hire more programmers, but how to structure career progression within a PK/PD programming function so that Domain 4 exposure is deliberately built rather than acquired incidentally. Organizations that rotate programmers across DDI, bioequivalence, and pediatric-extrapolation studies as a matter of career development, rather than assigning staff to whichever specialized study happens to need coverage at a given moment, are more likely to build the kind of cross-sub-area fluency this framework suggests is otherwise slow to develop. This has a direct bearing on succession planning as well: because Domain 4 expertise is comparatively concentrated among a small number of experienced staff at most organizations, losing even one or two such programmers can create a coverage gap that is considerably harder to backfill than a Domain 1 or Domain 2 vacancy, given how narrowly the relevant experience is distributed across the broader clinical programming labor market.

Table 2: Domain-specific training and staffing implications for regulatory-facing PK/PD programming teams.

Domain	Transferability From Adjacent Domain	Training/Staffing Implication
1. NCA	High foundational skill for all PK programming	Appropriate entry point for early-career programmers
2. Population PK Support	Moderate builds on Domain 1 but requires modeling-adjacent literacy	Benefits from cross-training with pharmacometrics team
3. Exposure-Response Support	Moderate requires endpoint-linkage judgment beyond dataset mechanics	Requires familiarity with therapeutic-area endpoints
4. Specialized Regulatory Analyses	Low DDI, bioequivalence, pediatric extrapolation each largely non-transferable from one another	Requires dedicated, sub-area-specific training track rather than general seniority

Limitations are worth stating plainly. This article synthesizes existing guidance and practice literature rather than presenting original empirical or benchmarking data. The ordinal intensity comparisons in Table 1, Table 2, Figure 1, and Figure 2 are illustrative

syntheses drawn from the literature reviewed, not measured process data, and are labeled as such wherever they appear. The framework is best read as a structured basis for training and staffing discussion, not a validated measurement instrument that has been tested against actual staffing outcomes or productivity metrics across multiple organizations. The bridging pharmacology literature drawn on here to frame the modeler-programmer relationship is itself an evolving field, and quantitative systems pharmacology approaches in particular continue to reshape what the modeling side of this handoff requires.¹⁴

A further limitation concerns generalizability across therapeutic areas and organizational contexts. The domain boundaries described here were derived from guidance documents and methodology literature that apply broadly across therapeutic areas, but the practical weight of each domain within a given development program varies considerably by indication. An oncology program with an adaptive dose-optimization design may place unusually heavy demand on Domain 3 exposure-response work relative to a program in a more established therapeutic area with well-characterized dosing, while a program developing a drug for a historically understudied pediatric indication may find Domain 4 pediatric extrapolation work dominating its programming resource needs in a way this general framework does not fully anticipate. Organizations applying this framework to their own resourcing decisions should treat the four domains as a starting structure to be calibrated against their own program mix, not as a fixed allocation formula. Finally, this framework was constructed from publicly available guidance and peer-reviewed literature rather than from direct observation of programming teams in practice, meaning the transferability judgments in Table 2 reflect an informed reading of the technical literature rather than a systematically collected account of how programmers themselves experience moving between domains. A study designed to directly survey or observe PK/PD programming staff across multiple organizations would be a natural next step for testing whether the escalation pattern described here matches lived staffing experience as closely as the literature comparison suggests.

CONCLUSION

Statistical programming support for PK/PD modeling splits into four domains, each with distinct technical conventions and an escalating share of regulatory consequence. Recognizing that division has practical value for training design, staffing models, and workforce planning at a moment when regulatory expectations for model-informed drug development keep expanding.

CREDIT AUTHOR STATEMENT

Radhika Aitha: Conceptualization, Methodology, Investigation, Writing – Original Draft, Writing – Review & Editing, Visualization.

ACKNOWLEDGMENTS

The authors used Claude (Anthropic, claude-sonnet model) to assist with manuscript drafting and language refinement. The authors take full responsibility for the accuracy, originality, and integrity of all content presented in this manuscript, and confirm that all data, analysis, and conclusions reflect their own independent work and verified sources.

Source of support: Nil; Conflict of interest: Nil.

REFERENCES

1. Food and Drug Administration. Population Pharmacokinetics: Guidance for Industry. Silver Spring (MD): FDA; 2022. Available from: <https://www.fda.gov/media/128793/download>
2. Federal Register. Population Pharmacokinetics; Guidance for Industry; Availability. 87 Fed Reg 6435 (Feb 4, 2022). Available from: <https://www.federalregister.gov/documents/2022/02/04/2022-02355/population-pharmacokinetics-guidance-for-industry-availability>
3. Food and Drug Administration. Pharmacokinetic-Based Criteria for Supporting Alternative Dosing Regimens. Silver Spring (MD): FDA. Available from: <https://www.fda.gov/media/151745/download>
4. Wu G, et al. A Retrospective Evaluation of Allometry, Population Pharmacokinetics, and Physiologically-Based Pharmacokinetics for Pediatric Dosing Using Clearance as a Surrogate. CPT Pharmacometrics Syst Pharmacol. 2019. Available from: <https://ascpt.onlinelibrary.wiley.com/doi/10.1002/psp4.12385>
5. Liu T, Ghafoori P, Gobburu JVS. Allometry Is a Reasonable Choice in Pediatric Drug Development. J Clin Pharmacol. 2017;57(4):469-475. Available from: <https://accpl.onlinelibrary.wiley.com/doi/abs/10.1002/jcph.831>
6. Mahmood I. An Update on the Use of Allometric and Other Scaling Methods to Scale Drug Clearance in Children: Towards Decision Tables. Expert Opin Drug Metab Toxicol. 2021;18(2). Available from: <https://www.tandfonline.com/doi/abs/10.1080/17425255.2021.2027907>
7. Food and Drug Administration. Clinical Drug Interaction Studies Cytochrome P450 Enzyme- and Transporter-Mediated Drug Interactions: Guidance for Industry. Federal Register (Jan 23, 2020). Available from: <https://www.federalregister.gov/documents/2020/01/23/2020-01064/clinical-drug-interaction-studies-cytochrome-p450-enzyme-and-transporter-mediated-drug-interactions>

8. Sudsakorn S, Bahadduri P, et al. 2020 FDA Drug-Drug Interaction Guidance: Comparison Analysis and Action Plan by Pharmaceutical Industrial Scientists. ResearchGate. Available from: <https://www.researchgate.net/publication/342341673>
9. Food and Drug Administration. Statistical Approaches to Establishing Bioequivalence: Guidance for Industry. Silver Spring (MD): FDA. Available from: <https://www.fda.gov/media/167453/download>
10. Food and Drug Administration. Statistical Procedures for Bioequivalence Studies Using a Standard Two-Treatment Crossover Design: Guidance for Industry. Silver Spring (MD): FDA.
11. Statistics in Biopharmaceutical Research. Methods for Non-Compartmental Pharmacokinetic Analysis With Observations Below the Limit of Quantification. 2019. Available from: <https://www.tandfonline.com/doi/full/10.1080/19466315.2019.1701546>
12. CDISC. ADaM popPK Implementation Guide v1.0. Clinical Data Interchange Standards Consortium. Available from: <https://www.cdisc.org/adam-poppk-implementation-guide-v1-0>
13. International Society of Pharmacometrics / Higher Logic. Basic Data Structure for ADaM PopPK Implementation Guide Version 1.0. Available from: https://higherlogicdownload.s3.amazonaws.com/ISOP/ab9a664e-5553-438f-b035-82b1e69e5009/UploadedImages/Documents/BDS_for_ADaM_PopPK_Implementation_Guide_v1_0_0.pdf
14. Kapitanov E, et al. Bridging the Gap: Integrating Quantitative Systems Pharmacology and Pharmacometrics in Drug Development. Clin Pharmacol Ther. 2026. Available from: <https://ascpt.onlinelibrary.wiley.com/doi/10.1002/cpt.70191>
15. Federal Register. Prescription Drug User Fee Act of 2023 VII Meetings Program for Model-Informed Drug Development Approaches. (Jan 11, 2023). Available from: <https://www.federalregister.gov/documents/2023/01/11/2023-00389/>
16. PHUSE / Wiki.cdisc.org. Implementation of CDISC ADaM in the Pharmacokinetics Department. Available from: <https://wiki.cdisc.org/download/attachments/55972337/4.Implementation%20of%20CDISC%20ADaM%20in%20PK%20dpt.pdf>
17. Jansen KM, et al. A Regression Approach to Visual Predictive Checks for Population Pharmacometric Models. CPT Pharmacometrics Syst Pharmacol. 2018. Available from: <https://ascpt.onlinelibrary.wiley.com/doi/full/10.1002/psp4.12319>
18. Food and Drug Administration. Exposure-Response Relationships Study Design, Data Analysis, and Regulatory Applications: Guidance for Industry. Silver Spring (MD): FDA. Available from: <https://www.fda.gov/media/71277/download>
19. Food and Drug Administration. Guidance for Industry: End-of-Phase 2A Meetings. Silver Spring (MD): FDA. Available from: <https://www.fda.gov/media/72211/download>
20. UDSpace / PMC. Model-Informed Drug Development: Addressing the Critical Need for Training in the Promising New Field. 2026. Available from: <https://pmc.ncbi.nlm.nih.gov/articles/PMC13048755/>
21. Applied Clinical Trials Online. Phase III Trial Failures: Costly, But Preventable. Available from: <https://www.appliedclinicaltrialsonline.com/view/phase-iii-trial-failures-costly-preventable>