

Teaching Statement

Habibolla Latifzadeh, Ph.D.

Stanford University — Department of Anesthesiology, Perioperative and Pain Medicine
Data Science Faculty Position

My teaching is shaped by two convictions formed during my own training. The first is that students retain a method only when they meet it from at least three directions—a derivation, a piece of working code, and a worked example on data they care about. The second is that mathematics and biomedical data analysis are most teachable when the instructor is honest about what each method assumes and where it fails. These convictions have guided every course I have led or assisted.

Record. At West Virginia University I served as Teaching Instructor for *Algebra with Applications* in Fall 2021, where I designed the weekly problem sets, wrote the exams, and built a small grading pipeline that re-ran student notebooks to flag silent failures—uninitialized variables, hidden state from prior cells—so that students learned to defend the correctness of a result rather than its appearance. I also taught and supported the quantitative-methods courses that anchor modern data science: *Applied Linear Algebra*, *Advanced Linear Algebra*, and *Numerical Analysis*. As founding President of the WVU SIAM Student Chapter (March 2022–present), I built a recurring student-led seminar and invited-speaker program bridging applied mathematics and data-driven research, which gave me a second venue—outside the classroom—to test how the same material lands across audiences with different preparation.

Philosophy. I teach mathematical and statistical methods as a craft. In Linear Algebra, the singular value decomposition is first a theorem about orthogonal change of basis, then a NumPy call, then a five-line image-compression demonstration. Students who only see one of those three forget the other two within a semester; students who see them aligned remember the geometry years later. In a clinical-machine-learning module I built during my NSF BridgesDH NRT Fellowship in AI and Machine Learning for Digital Health, I asked students to fit a logistic regression to a publicly available sepsis cohort and then identify three reasons their AUC was misleading. The point was not to trick the class; it was to teach calibration, subgroup performance, and label noise as part of the modeling task itself rather than as caveats appended at the end. This kind of honest treatment of assumptions matters even more in clinical settings, where the cost of a poorly calibrated model is concrete and immediate—and where clinicians cannot wait three semesters to apply what they learn. I plan to teach this way at Stanford, designing every course so that students leave with a working artifact—a fitted model, a validated cohort, a reproducible pipeline—they can use in their own clinical or research setting before the course ends.

Courses I would offer at Stanford. For undergraduates entering data science from outside mathematics, an introductory course on *Statistical Learning for Health Data*—penalized regression, tree ensembles, conditional-independence testing, and survival analysis—each method paired with a publicly available clinical dataset so students leave with a portfolio of reproducible analyses. For master’s students in data science or biomedical informatics, *Causal Inference for Observational Health Records*, covering targeted maximum likelihood, instrumental variables, and sensitivity analysis alongside the practical realities of EHR data: missingness mechanisms, irregular sampling, label leakage from documentation patterns, and governance constraints. For doctoral students, two seminars I would be glad to teach: *Bayesian Methods for Biomedical Networks*, drawing on my own research in ensemble structure learning and bootstrap-stability analysis for gene-regulatory and immunological networks; and *Optimal Transport for Biomedical Data Alignment*, covering entropic and unbalanced transport, with attention to relaxed-marginal formulations of the kind I have

developed for cross-species flow-cytometry mapping.

A third doctoral course I would propose is the flagship of this group: *Uncertainty-Aware Clinical Machine Learning for Perioperative and Critical-Care Data*, organized around three units—probabilistic calibration, distribution shift, and subgroup auditing—and around the data sources that actually shape clinical AI: electronic health records, laboratory values, medications, high-frequency physiologic waveforms, imaging-derived variables, clinical notes, and biomolecular profiles. Students would study supervised prediction, unsupervised phenotyping, representation learning, causal inference, and model auditing. The final project would require each student to define a cohort, justify inclusion and exclusion criteria, construct features, train and evaluate a model, and write a short report that combines quantitative results with clinical interpretation. The course is directly aligned with my postdoctoral work at Duke Biostatistics & Bioinformatics, my Aim 3 contribution on the NIH/NIAID R01 with the Permar Lab at Weill Cornell Medicine (in silico vaccine-efficacy modeling for congenital CMV), and my incoming Visiting Scholar & Co-Investigator role at the MIT Laboratory for Computational Physiology.

Across these courses I would use a layered assignment structure: a required core that every student completes, optional technical extensions for students seeking greater depth, and short interpretation prompts that require every student to explain results in plain language. This structure lets students contribute from their strengths while developing skills outside their initial training—particularly important for cohorts that mix quantitative, biomedical, and clinical backgrounds. Each course would also be built around a reproducibility expectation: students would submit code, data manifests, and a written analysis another student could rerun end-to-end on a SLURM cluster or a containerized workstation. Reproducibility is not separate from learning the methods; it is the test that the methods have been learned.

Mentorship. This summer I serve as Trainee Mentor and Hackathon Co-Organizer for the Multi-scale Immune Systems Modeling (MISM) Center of Excellence Summer Trainee Program at Duke, an NIH-funded immunology multiscale-modeling cohort. My contributions span project guidance, data-analysis support, code review, and hackathon coordination—a setting that combines mentored research, peer learning, and short-format intensive instruction in exactly the proportions a clinical-data-science training pipeline requires. The experience has sharpened my sense of how to scaffold trainees who arrive with strong domain intuition but uneven computational preparation, which is the same problem a Stanford perioperative-data-science cohort will present in a different vocabulary.

Earlier, I served as External Graduate Student Panelist for the NSF REU on *Discrete and Continuous Analysis in Appalachia* at Fairmont State University in 2022, and as invited graduate judge for the WVU Undergraduate Research Symposia in 2022 and 2023; together these roles gave me a closer view of how undergraduates from regional and Appalachian institutions decide whether graduate research is for them, and across both I saw the same pattern—strong technical execution paired with consistent underweighting of assumptions, limitations, and uncertainty—which I now address explicitly in every project rubric I write. At the graduate level, I mentored three M.Sc. students through thesis research at Shiraz University of Technology between 2010 and 2014, all of whom produced peer-reviewed publications.

The most useful thing I can then offer a research student is a clear standard of evidence combined with the removal of as many incidental barriers as possible: weekly meetings on a fixed schedule, draft feedback within two days, and an explicit distinction between what I am evaluating in their code, their writing, their presentation, and their choice of question. At Stanford, I would expect to

mentor across a wider population than I have so far: doctoral students in data science, statistics, biomedical informatics, and computer science; postdoctoral fellows building methods groups of their own; medical students, residents, and clinical fellows seeking quantitative training during research electives; and master’s students bridging into clinical research. The same standards apply across that range, but the framing changes. For a clinically trained mentee, I would build the methods around a question they bring from the bedside, and I would help them learn how to evaluate models, interpret uncertainty, and ask technical questions about validity. For a quantitatively trained mentee, I would build a clinical-context module around a method they want to develop, and I would help them learn how clinical data are generated and where measurement artifacts enter the pipeline. In both cases, I would expect them to replicate at least one figure from any paper they cite and to write down what they expected to see before running an experiment.

Inclusion. Many of my undergraduate students at WVU came from rural Appalachian backgrounds, and many were first-generation college students. I learned to write problem sets that did not assume cultural references or prior software exposure, to keep office hours that started on time and were genuinely useful, and to identify students who would never ask a question in front of peers but would happily ask one in a five-minute walk after class. The layered assignment structure described above is the same instrument used for a different purpose here: it lets students with weaker mathematical preparation engage seriously through the core and the interpretation prompts while quantitatively stronger students take the technical extensions. I would carry these practices to Stanford and apply them particularly to students entering data science from clinical and applied-science backgrounds whose mathematical preparation will vary—and to clinical trainees whose quantitative confidence often does not match their domain expertise.

Self-assessment. I have used end-of-course surveys, mid-term anonymous feedback forms, and exit conversations with research mentees to revise my own practice. After my first independent course, students told me my problem sets were well-structured but that my lecture pacing left too little room for questions; I now reserve the final fifteen minutes of every class for open discussion, and I have kept that practice since. I plan to continue collecting structured feedback at Stanford—short anonymous surveys during the quarter rather than only at the end, so that misunderstandings are corrected while the material is still being taught—and to share what I learn with departmental colleagues rather than keep teaching improvements private.

Teaching is the part of academic work where careful preparation produces the most visible returns for the most people. I would welcome the chance to bring that practice to Stanford, at the interface of Anesthesiology, Biomedical Data Science, and the broader data science initiative across the School of Medicine and engineering—training the next generation of clinicians, computational scientists, and clinician-scientists who will build the perioperative and critical-care analytics of the coming decade.

Habibolla Latifzadeh, Ph.D.