

A Conversation with Professor Weijie Su

Yuhua Zhu and Xiaowu Dai

Weijie Su is a Professor in the Department of Statistics and Data Science at the Wharton School of the University of Pennsylvania and, by courtesy, in the Departments of Computer and Information Science, Biostatistics, Epidemiology and Informatics, and Mathematics. He also serves as Co-Director of the Penn Research in Machine Learning Center. He received his Ph.D. in Statistics from Stanford University in 2016 and his bachelor's degree in Mathematics from Peking University in 2011.

His research interests span the statistical foundations of generative AI, high-dimensional statistics, privacy-preserving data analysis, and optimization. His work has contributed to a broad range of topics, including LLM watermarking, alignment and ranking, privacy-preserving data analysis, peer review mechanisms in machine learning, convex optimization, deep learning theory, and high-dimensional inference.

He is a founding Co-Editor of Statistical Learning and Data Science and serves as an Associate Editor for JASA, AOAS, JMLR, Foundations and Trends in Statistics, Harvard Data Science Review, and Operations Research. He currently serves on the Organizing Committee of ICML 2026 as Scientific Integrity Chair, where his isotonic mechanism will be deployed to enhance peer review.

His work has been recognized with many honors, including the Stanford Theodore Anderson Dissertation Award, NSF CAREER Award, Sloan Research Fellowship, IMS Peter Hall Prize, SIAM Early Career Prize in Data Science, ASA Noether Early Career Award, ICBS Frontiers of Science Award in Mathematics, IMS Medallion Lectureship, and the Outstanding Young Talent Award in the 2025 China Annual Review of Mathematics. He has authored two discussion papers in JRSSB and JASA and is a Fellow of the IMS and ASA. In 2026, Su received the COPSS Presidents' Award.

This interview was conducted by Yuhua Zhu and Xiaowu Dai from the Department of Statistics and Data Science at UCLA on April 21, 2026. The conversation covers Su's early life, his path into statistics, and his reflections on the future of statistics and artificial intelligence. The transcript of the interview is presented below.

Part I. Mathematical Beginnings and Early Life

Q: First of all, congratulations on receiving the 2026 COPSS Presidents' Award. Let us begin with your childhood. Could you tell us how you first became connected with mathematics?

Su: Thank you very much for this interview. I think it was around third or fourth grade in elementary school when I began to feel that I was reasonably good at mathematics. At that time, I was exposed to some word problems in mathematics. I found them fairly natural to solve, and I could solve them a bit faster than many of my classmates. It was probably during that period that I gradually realized I had some strength in mathematics, and my interest grew from there.

There is one experience that left a particular impression on me. It happened around fifth grade. Somehow, I came across a book on probability theory. I think it was probably a college-level textbook, though I no longer remember how it came into my hands. The book introduced Buffon's Needle experiment: if one draws parallel lines on a plane with spacing a , and randomly drops a needle of length l , then the probability that the needle intersects one of the lines is $p = 2l/(\pi a)$. If the experiment is repeated n times and the needle intersects a line k times, then one can estimate π by $2ln/(ka)$.

I actually did that experiment. Over a little more than a month, I must have dropped the needle tens of thousands of times, and my mother helped me with it too. In the end, I got something like 3.17, which was still a bit off from 3.14. Looking back, that was basically a Monte Carlo method in statistics.

Q: Besides classroom learning, were there other special experiences in your childhood, such as mathematical competitions or other extracurricular activities?

Su: In fact, I had not received systematic training in mathematical olympiads. My first participation in a math competition was quite accidental. In many schools in cities in Zhejiang, China, students

were organized specifically to take part in competitions, and many students would sign up. But in rural schools, very few people did this kind of thing, so there were not enough participants. Later, perhaps the local education authorities reminded schools that they should also organize students to participate, and the teacher asked in class who had reasonably good grades. That was how we were registered, more or less to give it a try.

I first participated in a math competition in eighth grade. I think I won third place at the city level and also received a second prize in Zhejiang Province. In that competition, most of the other award winners were ninth graders. Because of that competition, a middle school in the city invited me to transfer there and gave me a scholarship. What I remember particularly vividly is that the scholarship seemed to include a computer—that was the first computer I ever owned.

That was indeed an important turning point. It was the first time I moved from a rural school to a city school. Although the rural area and the city were not far apart—only a ten-minute bus ride or so—the school resources, the peer environment, and the overall atmosphere were very different. Without that math competition opportunity, I might not have been able to enter the best high school in the city through the regular entrance exam. And if I had not entered the best high school, my later trajectory might have been quite different.

Later, in high school, I continued to participate in math competitions. In my first year of high school, I took part in the Chinese Mathematical Olympiad and was offered admission to Tsinghua University. In my third year of high school, I again participated in the Chinese Mathematical Olympiad and later entered the National Training Team. In the end, I was admitted to Peking University.

Q: During the high school, did you already feel that mathematics was the academic direction that interested you most?

Su: Looking back, mathematics was probably the direction that interested me most at the time. But this also had to do with the environment in which I grew up. At that age, I did not have many concrete ideas about future career choices, nor had I developed other clearly defined interests. When I had just transferred to the city middle school, a teacher once handed out a form asking everyone to list their hobbies. I thought I should probably write down a hobby, but after thinking about it for a while, I really could not come up with anything

else. In the end, I wrote “mathematics.” In some sense, it was from that time that mathematics gradually became my most explicit interest.

I have often felt that mathematics is a very “egalitarian” subject. It does not require many material conditions, nor does it depend on expensive resources. As long as one has a book, anyone can begin to learn it.

Part II. From Peking University to Stanford

Q: During your four years at Peking University, what led you to choose statistics as your major?

Su: At that time, the School of Mathematical Sciences at Peking University had several major tracks. The first two years of coursework were more or less the same for everyone, and students began to specialize in the last two years. There was pure mathematics, or foundational mathematics; computational mathematics, which included partial differential equations and numerical analysis; and also statistics, financial mathematics, and information mathematics.

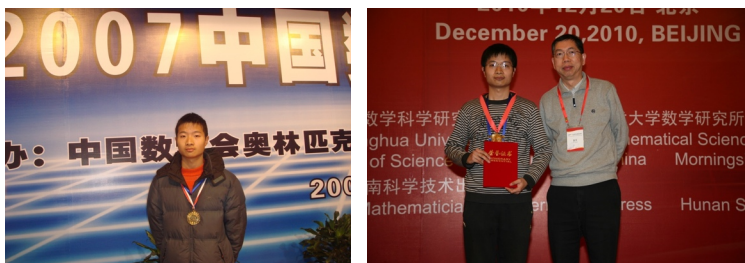
In fact, as an undergraduate, I did not study much statistics. I only took one statistics course. I studied more pure mathematics and computational mathematics, and I had some exposure to statistics. My later decision to apply for Ph.D. programs in statistics was not because I had a very clear intention from the beginning to move into statistics. It had a lot to do with the academic trend at the time. Around 2010 and 2011, compressed sensing was very popular—somewhat like AI and large language models today. You could see related work everywhere.

Later, I noticed that two representative figures in that field, Emmanuel Candès and David Donoho, were both in statistics departments. That had a great influence on me and became an important reason why I later applied to statistics Ph.D. programs.

Q: How was your Ph.D. application process?

Su: Overall, it was relatively smooth. Although I had not taken many statistics courses systematically as an undergraduate, my mathematical foundation was fairly solid, and my course grades were on the top.

There were also a few experiences that helped my application. One was that around the time I



(Left: Weijie at the Chinese Mathematical Olympiad Winter Camp in January 2007, where he won a gold medal and was selected for China's national training team.)

(Right: Weijie at the S.-T. Yau College Student Mathematics Contest in Beijing in December 2010, where he won a gold medal.)



Left: Weijie presenting the work on Nesterov's accelerated gradient method at the 28th NeurIPS conference in December 2014.

Middle: Weijie at Stanford University's Ph.D. commencement ceremony in June 2016.

Right: Weijie in Marseille, France, in June 2022, while attending a CIRM conference.



Left: Weijie at Enshi Grand Canyon in Hubei, China, in July 2024.

Middle: Weijie at the Paul A. M. Dirac statue at Florida State University in April 2026.

Right: Weijie in front of the Brown University gate in March 2026.

was preparing my Ph.D. applications in 2010, I participated in the S.-T. Yau College Student Mathematics Contest and won a gold medal. Around that period, there happened to be several related con-

ferences, and many scholars from overseas came to Beijing, including Jun Liu, Wing Wong, and Bin Yu. I took the opportunity to contact them and meet with them. Later, the application outcomes at

several schools were quite good.

At that time, Wing Wong was the chair of Stanford's statistics department, and I had met and spoken with him in Beijing. Later, Stanford admitted me to the Ph.D. program and also awarded me a Stanford Graduate Fellowship.

In addition, I had a little research experience. During the summer after my junior year, I interned at Microsoft Research Asia. The opportunity itself involved some chance. I no longer remember exactly how I obtained it, because most interns at Microsoft Research Asia came from computer science backgrounds, and there were many students from Peking University's School of Electronics Engineering and Computer Science and from Tsinghua's computer science program. The competition was very intense. I had a mathematics background, so I think I was lucky to get in. At Microsoft Research Asia, I worked on a project related to data and click behavior, somewhat like what we would now regard as basic recommendation or ranking problems. Today, it may not seem particularly complicated, but at that time it was quite cutting-edge and attractive. For an undergraduate, having the chance to engage with such problems was a very good opportunity, and it probably helped my Ph.D. applications as well.

Q: Did the real transition from mathematics to statistics happen during your Ph.D. in Stanford?

Su: Yes. The real transition mainly happened during my Ph.D. The first year at Stanford was basically devoted to taking courses seriously. Because I had not studied much statistics before, I needed to adjust to the statistical way of thinking. One very concrete challenge for me was using R for the first time. I had never used it before, so it took me some time to get used to it. That first year was mainly about coursework and building foundations. I did not have many other thoughts yet.

In my first year, I also took a course with Professor Emmanuel Candès, and he left a deep impression on me. He taught with great enthusiasm and energy. Although I later heard that he was gradually moving away from compressed sensing as his main focus, I admired him greatly at the time, so I eventually chose to work with him for my Ph.D.

The project that really helped me begin to un-

derstand statistical thinking was a project Emmanuel gave me on SLOPE. The problem was related to the Benjamini-Hochberg (BH) procedure in multiple testing. The BH procedure is a classical method for controlling the false discovery rate, and multiple testing itself is an important topic in statistical inference. The 2006 Annals of Statistics paper by Abramovich, Benjamini, Donoho, and Johnstone had already revealed a deep connection between the multiple testing and the high-dimensional sparse estimation. They showed that, by choosing thresholds adaptively in the spirit of multiple testing, one can achieve the minimax estimation rate when the sparsity level is unknown. We wanted to understand whether this connection could be extended from the relatively simple orthogonal model to the regression setting¹. It was through working on this problem that I learned many statistical ideas, especially the subtle but important differences between inference and estimation.

Previously, I had been more influenced by mathematical training, and I tended to think that a problem could ultimately be mathematized: once the definitions were clear and the theorem was proved, the problem seemed to be solved. But later I gradually came to understand that statistical inference is not entirely like that. For many inferential problems, the core lies in interpretation. For example, in what sense is a result significant, and when is it not significant? How should a discovery be interpreted? What is its relationship to the original scientific question? These issues are difficult to capture fully through a purely axiomatic language. In some sense, if one formalizes them completely, one may lose part of the interpretive space that made them meaningful in the first place.

This is something I learned a great deal about from Emmanuel. After obtaining a result, he would often ask: What does this result really say? How is it related to existing work? What is it useful for? What is the point? At first, I was not used to these questions. From a mathematical perspective, if a result is correct and the proof is valid, that might seem sufficient. But he cared more about what the result meant for the statistical problem itself and how it related to what we were really trying to understand or express.

Looking back now, that training was very important. It helped me gradually develop my own research taste. Good research is not only about de-

¹Bogdan, M., Van Den Berg, E., Sabatti, C., Su, W. and Candès, E.J., 2015. SLOPE—adaptive variable selection via convex optimization. *Annals of Applied Statistics*, 9(3), 1103-1140.

²Su, W. and Candès, E., 2016. SLOPE is adaptive to unknown sparsity and asymptotically minimax. *Annals of Statistics*, 43(3), 1038-1068.

ciding whether a statement is true, nor only about proving a theorem. More importantly, it is about understanding why the result is worth studying, how it relates to existing knowledge, and whether it genuinely answers an important question. This process cannot rely entirely on formal mathematical language. It also requires domain knowledge, statistical interpretation, and the researcher's own judgment.

Q: That is an insightful reflection. Apart from the SLOPE project, were there other interesting topics during your Ph.D.?

Su: Emmanuel had previously done a lot of work on fast algorithms, and those problems were closely connected to computational mathematics and applied mathematics. Therefore, in his group, optimization was a very natural research direction.

At that time, his research interests were very broad. They included applied mathematics, signal processing, and optimization, as well as statistics and information theory. The group was highly diverse in terms of research directions. Sometimes students would take turns presenting their work, and other students might not always be able to follow everything, because the problems and backgrounds were so different. But looking back now, it was precisely this diverse training that had a major influence on my later research style. It made me more accustomed to looking at problems from a broad and open perspective, and it gave me some understanding of how researchers in related fields think. For example, I would not say that I belong to optimization or signal processing, but I roughly understand how people in those areas approach problems. That perspective came to a large extent from the open, interdisciplinary, and diverse training in Emmanuel's group.

My later work on Nesterov's accelerated gradient method also arose in that context. That paper was joint work with Stephen Boyd and Emmanuel, and it used differential equations to characterize Nesterov's accelerated gradient method³. Looking back, it was a very fortunate project—perhaps one of the projects in which I invested relatively little time but gained a lot. Projects like that are rare in academic research. That is why I often feel that, in research, effort and judgment are of course important, but luck is also very important.

Q: Were there other faculty members who

had an important influence on you during your Ph.D.?

Su: Besides Emmanuel, several other faculty members at Stanford also had influence on me. Many of those influences did not necessarily come from direct collaboration, but more from classes and interactions.

For example, Persi Diaconis. I took his course and also heard some of his discussions with students. He gave me the impression of being an exceptionally scholar. He was already in his seventies at the time, but he still had a very strong curiosity about many problems. That left a deep impression on me. It is not easy for someone to continue to have genuine interest and curiosity after doing scholarship for so many years. So I feel that, as scholars, we should try to slow the decay of our curiosity as much as possible. That is something I hope I can do.

Wing Wong also influenced me. Although his research direction was not exactly the same as mine—he mainly worked in Bayesian and related areas—he was a top scholar with very high standards for research. I felt this more indirectly through his students. People would talk about how detailed he revised papers. That kind of rigor and persistence over many decades is very inspiring to younger scholars.

Iain Johnstone also left a deep impression on me. I took some of his courses. His lectures were very rigorous, logically clear, and technically solid. It is hard to say exactly how any particular lecture changed me, but this style of scholarship influences students gradually and subtly.

David Donoho's influence on me was more indirect. During my time at Stanford, I did not interact with him very much, perhaps only a few times. After graduation, at conferences or when he visited Wharton to give talks, we had more interactions. When I was at Stanford, everyone would say that his academic vision was extremely broad and that he understood statistics, applied mathematics, and computer science very deeply. A scholar like that sets a very high standard for students, even if only implicitly.

Lexing Ying was similar. He was a member of my dissertation committee and a scholar I admired. Although he does not work in statistics, he is very interested in statistical problems and often offers his own views. His curiosity-driven scholarly spirit was very infectious. Each time I spoke with him, I did not necessarily learn a specific technical point, but I was influenced by his interest in the problems

³Su, W., Boyd, S. and Candes, E.J., 2016. A differential equation for modeling Nesterov's accelerated gradient method: Theory and insights. *Journal of Machine Learning Research*, 17(153), 1-43.

themselves and by his curiosity.

Q: Many of your Stanford classmates later continued in academia. Did they also influence you?

Su: Yes. Many of my classmates from the Stanford cohort later continued in academia, including Bhaswar Bhattacharya, Yunjin Choi, Linxi Liu, Joshua Loftus, Stefan Wager, Chaojun Wang, Jing-shu Wang, Qingyuan Zhao, and others.

Looking back, that environment had a great impact on me. The influence of classmates is often more implicit and more random. Many times, it is hard to say exactly which classmate or which conversation influenced you, but the influence is real. At the time, everyone studied and discussed problems together, and naturally there was a kind of peer pressure. As a student, being surrounded by very strong peers could feel stressful. But after graduation, looking back, it was actually a great benefit. In such an environment, one's standards are raised almost unconsciously, and one gradually learns how to think about problems and how to do research.

After leaving graduate school, it is very hard to find again that kind of long-term, intensive environment in which people learn and exchange ideas together. So I think the influence of classmates is no less important than that of teachers; it is just more subtle. Looking back, whether during my undergraduate years or my Ph.D., being together with such outstanding classmates certainly helped me later pursue an academic career.

Part III. Finding the Intersection Between Statistics and AI

Q: After receiving your Ph.D. in 2016, you joined the Department of Statistics and Data Science at the Wharton School of the University of Pennsylvania and began your independent academic career. From doctoral training to gradually building your own research program, how did this process unfold?

Su: Looking back now, many of the choices in this process were not fully planned at the time. They emerged gradually through repeated explorations. Perhaps academic research is often like that. Many times, one does not know from the beginning where the most suitable direction lies. Instead, one adjusts along the way and gradually moves closer to

problems that one finds more important and more suitable. Only in retrospect does the trajectory appear clearer.

During my first two years at Wharton, I was mostly exploring. During my Ph.D., I had mainly worked in two directions: high-dimensional statistics and problems related to optimization. When I graduated in 2016 and joined Wharton, AlphaGo was attracting broad attention. AlphaGo was a landmark event. It made many people realize that machine learning and AI could achieve very powerful capabilities. At that stage, it was not easy to continue pushing forward the two directions I had worked on during my Ph.D. I wrote some papers, but the paper acceptance process was not smooth. So the first two years were indeed difficult. At that time, some popular directions in statistics, such as causal inference or single-cell data analysis in biostatistics, were not areas in which I had been trained during my Ph.D. I also thought about whether I should move into those directions, but later I felt it would be difficult to make a real transition because I lacked the relevant training and background. After some time, I began gradually moving toward machine learning. My idea was to see whether there were problems in machine learning that were genuinely statistical, such as problems involving inference, uncertainty, or reliability. In other words, I wanted to find problems that had not been sufficiently considered within machine learning but where statistics could play a role.

At first, this process was somewhat awkward. Because I came from a statistics background, when I first went to machine learning conferences, many people did not know me, and I did not have many friends there. But my attitude was to be humble: it did not matter if people did not know me, and it did not matter if I did not have friends there. I would simply go to posters and talks, trying to understand what people were doing and what problems they cared about.

It continued for some time. Gradually, I began to understand where the intersection between statistics and machine learning might lie. That intersection is not easy to find. On the one hand, statisticians have to feel that the problem is genuinely statistical. On the other hand, people in machine learning have to feel that the problem addresses something they care about.

Q: What were some key moments that made the intersection between statistics and machine learning clearer?

Su: Looking back now, an important turning point was indeed the emergence of ChatGPT. For me personally, it was a fortunate moment. Before that, I had already accumulated some research experience related to machine learning, including work on differential privacy⁴. Those works did not necessarily have an immediate impact at the time, but they helped me gradually become familiar with the questions and research language of the machine learning community.

After ChatGPT appeared at the end of 2022, the broader academic environment changed dramatically. Large language models made it clear to many people that artificial intelligence is not merely an internal problem of computer science, but something that will have a profound impact on many disciplines. The statistics community also began to pay much more attention to machine learning and artificial intelligence. In the past, it might have taken a lot of effort to explain why certain problems were important. But starting in 2023, people became more willing to discuss these issues actively and more interested in what role statistics could play.

This was certainly a major boost for my research. Many of the problems I had already been interested in—uncertainty, privacy, evaluation, and reliability—became more natural in the era of large language models and easier for others to understand. There also began to be more space for communication between statistics and machine learning around these topics.

Q: From your perspective, what is the most important role of statistics in the age of artificial intelligence? And could artificial intelligence, in turn, drive the development of statistics itself?

Su: This is something I am still thinking about. For example, alignment⁵, watermarking⁶, and reliability⁷ are directions I have worked on relatively extensively. These problems are important, but in many cases, they still involve bringing existing statistical ideas and tools into artificial intelligence: using a statistical perspective to understand large

models and using statistical methods to solve concrete problems in large models.

What is more challenging, and also more interesting to me, is whether artificial intelligence can in turn drive the development of statistics itself. In other words, will the new problems raised by the AI era lead statistics to develop new concepts, new methods, or even new intellectual frameworks? At the moment, I do not think we have a very clear answer. But I believe this is a very important question for the future. Statistics should not only treat artificial intelligence as a new application domain. It should also ask whether it can draw new intellectual resources from artificial intelligence and thereby further expand statistics itself.

Q: Looking ahead over the next few years, do you expect your research to continue focusing on large language models and generative artificial intelligence?

Su: In terms of concrete directions, I think I will continue to focus on large language models. Their impact on human work is enormous, and this impact is not limited to any single field. Language itself is a fundamental medium for human expression, reasoning, and collaboration. Now, with the development of agents, the scope of applications for large language models has become even broader, and they are influencing our work and lives in a very comprehensive way.

More fundamentally, I am interested in what foundational role statistics can play in the development of artificial intelligence. Today's artificial intelligence, especially generative AI, still has a strong probabilistic character. As long as generative AI remains an important direction, statistical ideas about randomness, uncertainty, and data will have great value.

Of course, some people believe that generative models alone may not be sufficient to reach artificial general intelligence, and that other paths or frameworks, such as world models, may be needed. But before those directions truly mature and show clear results, generative AI remains the most impor-

⁴Dong, J., Roth, A. and Su, W.J., 2022. Gaussian differential privacy. *Journal of the Royal Statistical Society Series B: Statistical Methodology*, 84(1), pp.3-37.

⁵Xiao, J., Li, Z., Xie, X., Getzen, E., Fang, C., Long, Q. and Su, W.J., 2025. On the algorithmic bias of aligning large language models with RLHF: Preference collapse and matching regularization. *Journal of the American Statistical Association*, 120(552), 2154-2164.

⁶Liu, K., Long, Q., Shi, Z., Su, W.J. and Xiao, J., 2025. Statistical impossibility and possibility of aligning llms with human preferences: From Condorcet paradox to Nash equilibrium. *Annals of Statistics*, to appear.

⁷Li, X., Ruan, F., Wang, H., Long, Q. and Su, W.J., 2025. A statistical framework of watermarks for large language models: Pivot, detection efficiency and optimal rules. *Annals of Statistics*, 53(1), 322-351.

⁸Su, W., 2026. Do large language models (really) need statistical foundations? *Annals of Applied Statistics*, 20(1), 724-743.

tant main thread. Precisely for that reason, I believe statistics still has much room to contribute in this era.

Part IV. The Responsibility of Statistics in the Age of AI

Q: From joining Wharton in 2016 to receiving the COPSS Presidents' Award in 2026, exactly ten years have passed. This is a remarkably rapid trajectory for a statistician. Looking back, what does this award mean to you?

Su: First of all, I feel very fortunate. There are many outstanding colleagues in statistics who are equally deserving of this kind of recognition; perhaps it is simply a matter of timing. So I am very grateful to receive this award, and I also feel that luck played a large role.

At the same time, I also feel some pressure. The accomplishments and standards of the senior scholars before me are very high. The award is certainly a recognition, but it also reminds me that I should not regard it merely as an outcome. I should regard it as a responsibility.

After receiving the award, I also felt some clear changes. For example, after the announcement, I received many more invitations to give talks, and my academic visibility certainly increased. This visibility is both an opportunity and a reminder. Especially against the backdrop of the rapid development of artificial intelligence, I feel that I need to think more seriously about what role statistics can play and how statisticians should participate in these new problems.

Therefore, after receiving the award, I also remind myself to hold my work to a higher standard. At the very least, I hope that the research I am doing can, to some extent, reflect the value of statistics and contribute a little to the dialogue between statistics and artificial intelligence. In that sense, one of the most obvious psychological changes after receiving the award is that, in addition to gratitude and a sense of luck, I also feel a greater sense of responsibility.

Q: For young students, what is uniquely attractive about statistics in the AI era?

Su: I think statistics is a very good direction, and I very much hope that more young people will

enter the field. At its core, statistics is about understanding the world through data, randomness, and uncertainty. As the world becomes increasingly datafied and digitized, we will face more and more data, and the uncertainty contained in that data will become increasingly complex. In this sense, statistics is a field with enduring vitality.

Data and uncertainty will not disappear. They will continue to arise in new scientific, technological, and social problems. So in the long run, statistics is not merely a short-term hot field. It is a relatively stable discipline that continually renews itself. For young people, it is a direction well worth investing in.

I also have a personal observation. Compared with many other fields, statisticians may be a relatively calm group of people. Perhaps this is because statistical training gives us a deeper understanding of randomness and uncertainty in the world. Many things naturally fluctuate, and outcomes cannot be completely controlled by individuals. After long-term exposure to this kind of training, one may become more accepting of uncertainty and develop a kind of natural immunity to external ups and downs.

So I think statistics is not only a promising academic field; it is also a way of looking at the world. It teaches people how to make judgments under uncertainty, how to find structure in noise, and how to face change with a more balanced mindset. From this perspective, if a young person hopes to enter a field that has long-term value and also trains one to be rational and steady, statistics is indeed a very worthy choice.

Q: What advice would you give to young students with a statistics background who hope to enter AI?

Su: I think the first thing is still to build a solid foundation in statistics. Of course, "building a solid foundation" sounds like a very ordinary statement, but it is genuinely important and does require time.

The environment facing young scholars and students today is very different from that faced by earlier generations of statisticians. In the past, many statisticians could spend a very long time building an extremely deep mathematical and statistical foundation. Today, under the impact of the rapid development of artificial intelligence, time and attention are much more fragmented, and it is difficult to train oneself entirely along the old path. Even so, if a student hopes to pursue long-term academic research in statistics, the basic statistical ideas and training cannot be missing. Core concepts such as

estimation, inference, and uncertainty must be genuinely understood. Otherwise, it becomes unclear where one's disciplinary identity lies and what the unique contribution of statistics is.

On that basis, if one is interested in AI, I think it is very much worth paying attention to. AI is not an ordinary application area. In some sense, it is a large-scale attempt to mechanize human intelligence. It is trying to transform many tasks that were originally performed by biological intelligence into tasks that machines can perform. At present, it has already replaced or assisted human intelligence in many respects, and its upper limit has not yet been fully revealed.

This will also raise many fundamental questions. For example, if AI can approach or even exceed human intelligence in more and more intellectual tasks, how will academic research be conducted in the future? Where will the agency of human researchers lie? Every discipline needs to confront these questions, and statistics is no exception.

More importantly, the current development of AI is deeply connected to statistical thinking. Generative models fundamentally rely on random generation, probabilistic modeling, and the representation of uncertainty. Although these techniques were not necessarily invented directly by statisticians, they do contain many core statistical ideas. One might say that, in some sense, AI is extending statistical thinking to new settings on a massive scale.

Of course, there are also deeper questions here. For example, is randomness essential to intelligence, or is it merely a feature of the current technological implementation? Can intelligence ultimately be made fully deterministic? These questions may sound somewhat philosophical, but statistics is capable of participating in them. After all, statistics has long dealt precisely with the relationship among uncertainty, data, and inference.

So I believe that, in the development of artificial intelligence, statistics will always have its own place. It cannot simply be replaced by mathematics, nor can it be completely replaced by computer science. Statistics has its own distinctive ideas and value. For young students, the key is to preserve the core training of statistics while also actively engaging with the new problems raised by artificial intelligence. Only then will it be possible to find contributions that truly belong to statistics in this new era.

Q: AI has led to a rapid increase in the number of papers, and peer review is under growing pressure. You have studied mechanism designs in peer review and have also participated in related practice⁹. How do you view peer review in the age of AI?

Su: Peer review is indeed a very important issue. A few years ago, I began thinking about mechanism design problems in paper ranking and peer review, and at the time I proposed some relatively simple ideas. I did not expect that, several years later, these ideas would actually have the opportunity to be used in real conferences.

If we look at the issue more broadly, the challenges facing peer review are very clear. The entire academic community is discussing whether peer review is approaching a state of overload. On the one hand, the number of submissions has increased sharply. In particular, with the help of AI tools, the cost of writing and presenting research has been greatly reduced, so the number of papers naturally grows rapidly. On the other hand, researchers' time is becoming increasingly fragmented, and the time they can devote to careful reviewing is actually decreasing.

More importantly, AI is also changing some of the traditional signals we use to judge paper quality. In the past, if a paper was written very clearly and in excellent English, that often suggested that the authors had invested a great deal of time and that the quality of the paper was probably not too poor. But after the emergence of AI writing tools, this judgment is no longer reliable. Even if a paper's ideas are immature or its results are not solid, AI polishing can make it read very smoothly. Therefore, some of the intuitions and signals that traditional peer review used to rely on are becoming less reliable.

I think this is not a problem for statistics alone, but a problem the entire academic community must face. Peer review will surely reach some new equilibrium in the future. But what that equilibrium will look like, and how we can make it healthier and more effective, are questions very much worth discussing.

In this process, statistics should be able to play a role. Peer review can essentially be viewed as an estimation problem with a great deal of noise: we hope to estimate the true quality of a paper from a limited number of reviews that are noisy and may even be biased. This involves ranking, aggregation

⁹Su, B., Zhang, J., Collina, N., Yan, Y., Li, D., Cho, K., Fan, J., Roth, A. and Su, W., 2025. The ICML 2023 ranking experiment: Examining author self-assessment in ML/AI peer review. *Journal of the American Statistical Association*, pp.1-12.

of information, bias correction, uncertainty assessment, mechanism design, and many other issues. These are all areas where statistics can participate and contribute.

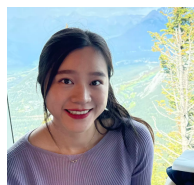
Of course, some people may say that, in the AI era, traditional peer review may no longer be needed. Papers can simply be posted online, and their impact can be determined by readers, citations, GitHub stars, or community feedback. But I do not think we can rely entirely on that approach. Without some form of quality-control mechanism, a large number of low-quality or even erroneous papers will enter the knowledge system. More importantly, a few years later, these papers may become training material for large language models. If they contain many incorrect conclusions, they will further contaminate future knowledge production.

So I think the importance of peer review will not diminish; it will simply become more difficult. Precisely because it is more difficult, we need to think seriously about how to improve it rather than simply give up on it. For statistics, this is a very natural and important problem. We have long dealt with noisy data, inference under limited information, and

uncertainty assessment. Peer review in the AI era is, at its core, about how to better estimate scholarly quality from complex and noisy information. Being able to contribute to this problem is, I believe, an important way in which statistics can contribute to the broader scientific community.

Q: Great! Thank you very much for taking your time to share your thoughts.

Su: Thank you very much for interviewing me.



*Yuhua Zhu, PhD,
Assistant Professor
the Department of Statistics and
Data Science,
UCLA.*



*Xiaowu Dai PhD,
Assistant Professor
the Department of Statistics and
Data Science,
UCLA.*

ICSA Mentor-Mentee Program: Introduction and Progress Update

Yushi Liu, Hanfei Xu, Jun Zhao, Xun Chen, Wei Shen, Hongzhe Li, Hongyu Zhao and Hongtu Zhu, Yabing Mai and Tsui-Feng Wu

Building on the successful launch of the ICSA Mentoring Lounge in October 2025, the ICSA Mentor-Mentee Program (IM²) has continued to grow its monthly, sector-focused Zoom sessions while expanding into its first in-person mentoring event. The program remains committed to its core mission of better preparing junior statisticians for the transition into industry by drawing on the expertise of seasoned professionals, and the sessions described below reflect both a broadening of the industry sectors represented and a deepening of community engagement.

Over the following months, the Mentoring Lounge expanded into three new industry sectors. The second session, on December 5, 2025, focused on the technology industry with mentors Chang Liu (Google) and Zerui Zhang (LinkedIn), drawing around 15 attendees out of 55 registrants. The third session, on February 20, 2026, was co-sponsored with the ASA Biopharmaceutical Section and centered on data science in pharma, featuring Haoda Fu (Amgen) and Zhaoling Meng (Sanofi) with roughly 15 attendees out of 55 registrants. The fourth session, on April 20, 2026, explored the path from statistics to the finance industry with mentors Lei Chen and Roy Tang, drawing 8 attendees out of 28 registrants. Across the three sessions, mentees raised questions on transitioning into each sector, the skills and curricula that matter most, the role