



Interactive
Comment

Interactive comment on “Water droplet calibration of a cloud droplet probe and in-flight performance in liquid, ice and mixed-phase clouds during ARCPAC” by S. Lance et al.

P. Chuang

pchuang@pmc.ucsc.edu

Received and published: 9 September 2010

Manuscript Review

Authors: S. Lance, C. A. Brock, D. Rogers, and J. A. Gordon

Title: Water droplet calibration of a cloud droplet probe and in-flight performance in liquid, ice and mixed-phase clouds during ARCPAC

Reviewer: Patrick Chuang, UC Santa Cruz

Summary:

C1394

Full Screen / Esc

Printer-friendly Version

Interactive Discussion

Discussion Paper



The authors describe a calibration system useful for in situ cloud drop probes and use it to carefully calibrate the DMT CDP. They find from observations that there are strong concentration-dependent biases in LWC, which they propose is due to coincidence. They use models to explore the artifacts due to coincidence and show that they can reasonably reproduce the observed biases; if they anything, the observations likely exhibit even stronger biases than reflected in the simulations performed.

General comments:

I think the work done here is of very high quality and of great importance to the field. There's no question that it's worthy of publication. That said, I'll have to admit that reading it was more frustrating than I think it should have been. One of the main issues seem to be that many of the key figures aren't as clearly explained as they should be. Details below. Interpreting the figures was also difficult for me sometimes, as I wasn't really sure what I was supposed to compare on a given figure. The figure captions are very sparse and could be a lot more helpful to the reader. But more than this, I think there are parts where the text is lacks explanation and clarity. I've tried to point out the main spots below. Definition of symbols in some places also lacking (especially n_D : you MUST be consistent with the usage of n_D !).

A second big issue is the author's use of "ideal CDP" instrument. The comments below dig into this issue somewhat, but I'll summarize here: the calibration system "calibrates" the instrument at maximum signal. The authors then conclude that drops passing through less intense parts of the beam leads to "undersizing" since the pulse amplitude is smaller. The result is a 2 μm shift between "ideal" and "realistic" CDP instruments. I'm not going to force the authors to change this interpretation, but I disagree with this thinking. I'm almost positive CDP utilizes a laser that has a Gaussian intensity profile (since most lasers exhibit this behavior) and without some very substantial modifications to create a homogeneous beam (which the CDP doesn't have), this is the instrument you're stuck with. So I think an "ideal CDP" still exhibits a transfer function that is not a delta function, i.e. a monodisperse drop distribution going through

C1395

Full Screen / Esc

Printer-friendly Version

Interactive Discussion

Discussion Paper



the CDP will lead to a range of pulse amplitudes. The obvious analog (to me, and which I know the authors will understand) is between treating an ideal DMA as one having delta function vs. triangular transfer functions. You might want delta functions, but the best you're going to get is a triangle with finite width. Treating the instrument as not having this is, to me, inappropriate. I really do think this transfer function approach would be more useful, too, as the transfer function width gives important information about the instrument resolution, and knowing the transfer function might also be used as the basis for an inversion to better recreate the parent or true distribution (as in DMA inversions).

The comments below don't necessarily reflect this, but I think the importance that the qualifier and sizer signals take on later in the paper require a lot more careful explanation of how this works in the instrument. A figure illustrating this would go a long way to helping the reader out.

Specific comments:

Note: sorry about the formatting/capitalization inconsistencies. Hopefully the comments are clear nevertheless.

page 3136, line 1: I prefer *in situ* in italics since it's Latin. NY Times agrees. You may not, however.

line 4: does this include sources of uncertainty? (random, not biases)

5: i'm not sure drop counting uncertainty is the same as view volume uncertainty. i'd make them separate.

5: "greater uncertainty": not necessarily - if uncertainty is bigger for small drops than big drops, then higher order products should be better off. yes, it's a bit pathological but possible...

14: at some point might want to point out this sample volume is drop size dependent for many instruments, even if it is treated as a constant.

20-23: this sentence seems a bit obtuse. it needs to make its point more directly, i think.

32: “standardized... sample area” that sounds a bit odd... like there’s a tube goes right up to the laser.

3137, 1: this doesn’t seem like a very useful sentence. i’d take the next sentence and divide into two: 1. response depends on ref index then 2. thus response to water must be calculated...

22: it’s not the flight that matters, though, is it? it’s more that it’s a population measurement at high speed rather than drop-by-drop at low speed, right?

32: “continuously... transmitting”: a bit roundabout - don’t you really mean that it’s difficult to consistently generate a population of known concentration and size distribution (and, for ice crystals, crystal shape).

3141, 5-6: I’m still a bit unclear how this works. This sentence makes it sound like the total light scattered is then divided between the qualifier and sizer, whereas later I infer that at least the intensity as a function of scattering angle information must be preserved. I think a picture would help this out immensely. I also think it’s such an important part of what happens later that you should take the time to explain this well.

10: Does the qualifier algorithm depend on drop size? If not, should it? Seems to me like it should (can’t use the same intensity threshold for 4 and 40 micron drops). If it does, does it make the sample area drop size dependent?

I recommend you add approximate dimensions (lateral and longitudinal) for both the SAQ and SAE to give people a sense for the values. I was surprised to see the difference between these in Fig 7. Are the qualified values substantially different from that for the FSSP?

24: “diffracted” why just diffraction being considered in this sentence? i think "scatter-ing" is better.

Full Screen / Esc

Printer-friendly Version

Interactive Discussion

Discussion Paper



24: "proportional..." not true for all drop sizes - only geometric scattering, definitely not Mie or Rayleigh scattering

3142, 2: and as you point out earlier, a model of the optical setup.

3145, 26: might this velocity change cause a temperature of the air? if so, any chance the drop size responds to this?

3147, 30: "10% greater" any sense for the uncertainty in this correction?

3148, 8: what is uncertainty in drop velocity measurements using this method?

3149, 5: Is "Fundamental" needed in the section title?

7: "initially calibrated" at what velocity?

10: this phrase is a bit confusing. does a "sizer" have a "pulse amplitude"?

i found this section very confusing. i read this paragraph at least six times and still can't totally understand it. the topic sentence talks about psl and glass beads. this third sentence doesn't seem to follow from the first two (nor does the fourth). then discussion goes back to the calibrations.

Leads to Fig 2 issues as well - it's not clear what is supposed to agree with what (like the red symbols, the way the cal was done, shouldn't actually correspond to the blue "curve", for example). left and right axes reversed in caption? also, if units are Watts, add to axis label. I like the uncertainties for the calibration points, but they're hard to see - I'd suggest making them darker (esp. for the PSL).

30: I'm also very surprised by this - and I think it deserves more discussion than this. Other people have been able to reproduce this behavior in other optical instruments, and I'm assuming it's been done for the FSSP though I don't know for sure. But in any case, it must say something important about either the optical model used for the instrument or the calibration method or the instrument itself. Is there some way to address this further?

Full Screen / Esc

Printer-friendly Version

Interactive Discussion

Discussion Paper



3150, 1-2: This is confusing. It kind of sounds like you're claiming that the red dots on Fig 2 represent smaller sizes for given pulse amplitude than the green or grey curve, but clearly you mean red dots vs blue curve. The point, I think, is that the blue curve is the green + grey curves shifted using some model. So the problem doesn't *have* to lie with the green/grey curve- it could be the model, right? And the model might be flawed since no Mie resonances show up in the red dots?

3: Is the cal constrained only to the center of the DOF (depth being along the laser path) or is it also constrained to the middle of the beam, meaning passing from 3 o'clock to 9 o'clock vs from 2 o'clock to 10 o'clock? I'm not 100% sure of the terminology, but I think of DOF as the former only.

What I'm confused by is why the calibration was done like this in the first place. Since the drops pass through lower intensity parts of the beam, isn't it necessary to calibrate based on all trajectories that will be accepted by the CDP, and not just the ones with max amplitude? Why should this 2 μm shift be needed? Couldn't one just traverse the cal drops through the beam until they're rejected by the instrument and then weight each stream by its likelihood?

8: In the first sentence where you introduce Fig 3 (above), you say "no averaging was performed and each data point represents a single droplet..." which is not consistent with the idea that Fig 3 represents D_v , which is a population measure.

Again, the issue of why one expects Fig 3 standard measurement data to line up along the 1:1 curve is odd - you selected only a small subset of possible trajectories and thus there's no reason to think that the CDP standard measurement is valid here. It only works as an average if you send drops along all trajectories, right?

23: It seems like the assumption here (and one that would clarify some of the above remarks) is that the ideal CDP has a transfer function that is a delta function. Obviously it's not. Although this would certainly completely change the way this paper is structured, I would suggest changing it so that you describe the transfer function instead.

Full Screen / Esc

Printer-friendly Version

Interactive Discussion

Discussion Paper



The valuable part about this is that one could then use an inversion routine (in the same way that DMA data is inverted, say) in order to account for the transfer function. One could fit the transfer function to different functions (lognormals? gamma functions?) and then describe the parameters of this with size.

If you choose not to go this route, then I think you need to be upfront about your assumption of the transfer function, and noting that no one, obviously, expects it to look that way.

In any case, the breadth of the transfer function (whether you ignore it or not) could be useful since it is another constraint on the size resolution of the instrument. Is there a regime where these sizing uncertainties dominate relative to Mie resonance uncertainties?

3151, 6: Can't you be more quantitative? Can you take the Fig 4 data and, assuming randomly distributed drop trajectories, come up with a good estimate for the count rate uncertainty? That would certainly be helpful since it directly affects concentration! Even cooler would be to do it as a function of drop size...

3152, 21-22: Move this sentence ("The ice-only... probe.") to previous paragraph.

26-27: Isn't the real test whether or not you see small particles at concentrations on the order of 10 to 100 times that of the precip-sized particles (as in Alexei Korolev's recent measurements)? So if you have 2 per liter of large particles, then you might expect shattering artifact concs on the order of 0.02 to 0.2 per cm^3 , right? And this is well within what was observed for concentration, right? So I'm not convinced this is a good case for testing shattering. Am I wrong here?

3153, 20: Can you explain the implications for choosing an inter-arrival distribution that's uniform and limited to $2 * \tau$? In the absence of any small-scale clustering, one would expect the distribution to be Poisson with mean (and standard dev) of τ , which means there's a long tail to larger τ values than you use here. Would using a more

Full Screen / Esc

Printer-friendly Version

Interactive Discussion

Discussion Paper



realistic distribution change the results?

24: this notation is inconsistent with Fig 6, where n_D is a concentration rather than a count rate. make it consistent!

3154, 1: I think this paragraph deserves a much better topic sentence!

28: I think Fig 8 needs to be explained better. I think "sizer sum" and "qualifier sum" are the total scattered light from all drops to each of these detectors. "Sizer signal" is, I'm assuming, the signal just from the drop of interest. These aren't labelled anywhere.

30: "inhomogeneous instrument response": i won't preach any further on this, but i'm not terribly fond of using this terminology - inhomogeneous response is built in to the design of this instrument and there's no way around it unless the optics are changed dramatically. i guess i'd view a "perfect CDP" as one that works as well as can be expected within the hardware design parameters.

to be less ambiguous, maybe you could refer to this as an instrument with a delta function as its transfer function? i didn't pick up what was meant by "inhomogeneous instrument response" on first reading.

3155, 8-9: "scatters additional light into the sizer" why isn't this additional light seen by the qualifier?

12-13 "droplets are undercounted": This sounds funny to me (the only way to undercount a single event is to make it at zero - strictly it's undercounting but it's really missing it entirely!). Maybe "undercounting of droplets occurs"?

25: I think you need to better explain Fig 9. Here are some questions I have about it that seem unexplained (the caption is exceptionally short and the text only helps a bit):

- Is this a monodisperse distribution or is there some width to it? If former, why is D_v used in the text (since D_v suggests other types of diameters, but if monodisperse, they're all the same)? If the latter, why isn't the distribution width described?

Full Screen / Esc

Printer-friendly Version

Interactive Discussion

Discussion Paper



- What is this pulse width associated with? The size or qualifier? Why is it permitted to be a free parameter? Isn't it constrained by the optics of the instrument? [OK, this is clarified in the next page by saying it's treated as independent: so why is it allowed to be independent when it's not? What's gained in doing so, and how would a more realistic simulation be different from these idealized ones?]

- What does "sub-100m variability in n_D " refer to? What is this variability? (It doesn't help that I don't know which n_D you mean as you use it for two different parameters). [explained much much later - I suggest removing this as described in more detail later]

- Apart from all the questions, what hurts my head the most is seeing (a) different kinds of simulations (homogeneous and inhomogeneous) and (b) simulations where both drop size and pulse width vary among the curves (oh, and having some of them fit to a line and others not). I don't really know how to interpret the figure - the text focuses on the change in drop sizes but pulse width changes and not in any obvious way (like smaller drops have smaller pulse widths).

This figure (9a and 9b) is also too small - I had to blow it up to 200% to make it comfortably legible. My eyes aren't great but I don't think they're that bad!

29: "The oversizing bias.." Are these supposed to be quantitative, i.e. representative of oversizing in the actual instrument, or qualitative, i.e. showing the approx. magnitude of the error and its variability with drop size? Again, I'm not sure what to take away from these simulations.

3156, 6: "Figure 10" You don't mention at all that you have plotted the prescribed and simulated distributions on totally different y-axis scales. I assume this is because of the undercounting issue. Why aren't they plotted on exactly the same scale to show the magnitude of this issue as well? Or if you keep them on different scales, at least make it clear. In either case, it seems worthy of discussion, no?

And in constructing this, what is the assumption about pulse width? In Fig 9 it's an

[Full Screen / Esc](#)[Printer-friendly Version](#)[Interactive Discussion](#)[Discussion Paper](#)

independent parameter. Presumably something was assumed here? Or not?

10: “variable... SAQ”: This is again, I believe, specifically referring to the fact that the drops see varying intensities in the sizing region, right? I find this sentence doesn't really say this - it just says the response varies but seems to lack a specific explanation. If I can offer a suggestion: can you come up with a specific term for this effect? Explain it carefully once and use it for the rest of the paper. For example, you could call it the “non-uniform illumination effect”.

17: Fig 11: The shaded area is, I assume, the same shaded area from Fig 5? You don't make it clear.

Is this Figure really a fair comparison? The simulations choose very specific values of size and the (to me, confusing) pulse width. Are these consistent with the data used for Fig 5?

3157, 2: Figure 12: I think if this were placed earlier (i.e. before Fig 9) it would have helped convince me that the simulations were representative of realistic conditions.

11: “We ran additional simulations...” I don't get why these simulations don't match any of the homogeneous ($L = 100$ m) simulations. They seem random.

If the only point here is that smushing all the drops into a smaller time period instead of distributing them homogeneously over 1s causes more coincidence, I'd suggest to remove these points. They confused me terribly when I was looking at Fig 9, and the point is rather obvious.

Figure 6: Define n_D .

Interactive comment on Atmos. Meas. Tech. Discuss., 3, 3133, 2010.

Full Screen / Esc

Printer-friendly Version

Interactive Discussion

Discussion Paper

