

Reply to Reviewer Comments to “Improving Models to Predict Holocellulose and Klason Lignin Contents for Peat Soil Organic Matter with Mid Infrared Spectra” (soil-2022-27)

Henning Teickner, Klaus-Holger Knorr

2022-10-26

Contents

Reply to Comments of Reviewer 1	1
General comments	1
Specific comments	2
Reply to Comments of Reviewer 2	3
General Comments	3
Specific Comments	4
Supplementary Information	13
Additional changes	14
References	15

Reviewer comments have prefix “Q” and our answers prefix “A”. “old” in combination with line numbers refers to lines in the previous version of the manuscript. “new” in combination with line numbers refers to lines in the updated version of the manuscript.

Reply to Comments of Reviewer 1

General comments

Q1: “This is an excellent manuscript based on an excellent rationale with sound research questions. It is very well structured and very well written. Building on previous data and models, Teickner & Knorr provide detailed explanations of the steps followed to evaluate the models, describing the limitations (representative of the training data sets, validation of the models, availability of key data of the models’ output, biases and uncertainties, etc.) and potential improvements. This research also shows the key importance of a detailed analysis of the models’ residuals as a way to identify model’s deficiencies. It is also a key finding that OM composition of mineral soils can also be accurately modeled. This research also shows the potential of MIR for SOM characterization and modelling and the importance of making data and code available for further improvements.

Regarding the general issue of the preselection of peaks before modelling, one possible alternative is to perform PCA on all detected peaks to reduce the dimension and extract the most relevant spectral signals. PCA is usually efficient in allocating into different components the overlapping effects of more than one compound on a given spectral band/region.”

A1: Thank you for this encouraging comment. It is true that there are many more modeling approaches which probably would have a similar predictive accuracy as the approach we used here (especially since the

sample size is small). For example, reviewer 2 suggested partial least squares regression (PLSR) as yet another alternative approach and there are many more approaches one could have tried (e.g. supervised PCR, iterative supervised PCR, interval PLSR, moving window interval PLS, etc; see e.g. Xiaobo et al. (2010)). We have the impression that PLSR is more commonly used than PCR since PLSR is more efficient in identifying latent variables related to the target variable being predicted (since it maximizes also the covariance of the computed latent variables to the target variable, whereas PCR only maximizes the variance of the computed latent variables, which comprise only the spectral data). That said, we think that it is good that different tools exist and, depending on the problem, one or the other approach may result in a better predictive accuracy. PCR (or a different dimension reduction approach) may be particularly useful if more data becomes available so that the computational expenses of our approach get too heavy (Piironen, Paasiniemi, and Vehtari 2020).

One advantage of the approach we used in comparison to dimension reduction approaches (e.g. PCR, PLSR), which summarize the original predictors into latent variables, is that coefficients for individual predictors are estimated more independently, since no latent variables are computed as summary of all predictors. Therefore, multiple regression with regularizing priors has the advantage of facilitating model interpretation (see e.g. Fig. 3).

We think that it is good to at least briefly mention these aspects to provide readers a better orientation. To this end, we added the following paragraph (old: l. 175; new: l. 185): “An alternative, popular, approach to approach 2 would be dimension reduction, for example via partial least squares regression, principal component regression, or variants of these (Xiaobo et al. 2010). In general, there are many alternative approaches which could be tested to use more information contained within the spectra than the original models, and many of these probably would result in similar predictive performances as approach 2 (regularization), especially when sample sizes are small (Xiaobo et al. 2010; Teickner, Gao, and Knorr 2022). An advantage of regularization is that model coefficients are estimated more independently than in dimension reduction approaches, which makes it more straightforward to interpret model coefficients. The key is that the approaches we chose are suitable to analyze our research questions.”

Specific comments

Q2: “L277-279. Positive and negative coefficients for peaks at 1150 and 1270 cm^{-1} (aromatic CH bending). This may be due to they reflecting different lignin structures: 1270 cm^{-1} surely corresponds to guaiacol (G) moieties and 1150 cm^{-1} could correspond to syringol (S) moieties (although S most representative band is at 1310 cm^{-1}). Both differ in content in the source peat vegetation and also show differences in resistance to degradation in peat depending on oxygen availability (see for example, Schellekens et al. 2012 Soil Biology and Biochemistry 53, 32-42). G lignin is less prone to degradation than S lignin and the G/S ratio of fen peat (more decomposed) is larger than that of bog peat (less decomposed) (see for example Martínez Cortizas et al., 2021 Boreas 50, 1161-1178).”

A2: Thank you for this information, this is indeed interesting. Building on this: The G/S ratio is also different in the training dataset sample types. Softwood samples and needle samples have a higher G content than hardwood samples and leaves and grasses (De la Cruz, Osborne, and Barlaz 2016). They also have a (slightly) lower Klason lignin content than hardwood samples and leaves and grasses (supplementary Fig. S4).

This indicates that since the bin at 1150 cm^{-1} is caused to some extent by S, and S has a larger abundance in samples with smaller (than average) Klason lignin contents, this may explain why the coefficient for this bin is negative. Likewise, the bin at 1270 cm^{-1} is caused to some extent by G, G has a larger abundance in samples with larger (than average) Klason lignin content, and this may explain why the coefficient for this bin is positive.

To consider this hypothesis in the main text, we changed the text to (old: l. 277 to 279; new: l. 294): “A plausible explanation for the negative sign of the coefficient for the 1150 cm^{-1} bin is that absorbance in this range is partly caused by syringyl units (S) (Kubo and Kadla 2005) and that training samples with higher S content have smaller Klason lignin contents (supplementary Fig. 4, De la Cruz, Osborne, and Barlaz (2016)), making this bin indicative of smaller Klason lignin contents. Likewise, a plausible explanation for the positive sign of the coefficient for the 1270 cm^{-1} bin is that absorbance in this range is caused predominantly by guaiacyl units (G) (Kubo and Kadla 2005) and that training samples with higher G content have smaller

Klason lignin contents (supplementary Fig. 4, De la Cruz, Osborne, and Barlaz (2016)), making this bin indicative of larger Klason lignin contents.”

Q3: “L424 I find this “implication” to be quite important and other ways of spectral normalization should be tested in the future.”

A3: We completely agree with this statement. We currently are not aware of a solution to this problem, other than to specify precisely under which conditions a prediction model produces erroneous predictions.

Reply to Comments of Reviewer 2

General Comments

Q1: “Several previous studies have used mid infra-red spectra to assess peat organic matter chemistry, particularly holocellulose and lignin content. One in particular had calibrated the spectral data with a set of chemical analyses using various (non-peat) organic materials. This paper aimed to re-examine the calibration using more refined models and gauge the suitability for applying them to peat and peat vegetation. The authors do a thorough job in applying some sophisticated statistical techniques to the data used by Hodgkins et al. (2018) and show how the calibration models can be improved and what the limitations of the current dataset are. From this point of view, the paper is useful and adds to the precision and accuracy of using infra-red spectra to rapidly assess peat chemistry, with possible application to other organic matter types. A weakness is that the authors have restricted themselves to the training chemical data from the Hodgkins et al. paper which are a curious set of paper, wood and plant leaves. In fact, a prior referee of the Hodgkins et al. paper described the set as being “bizarre” and actually it is surprising that such reasonable results are obtained from such an unlikely dataset. One would normally expect the training set for peat analysis to be based upon samples of actual peat, from as wide a set of sources as possible. Clearly, the authors have limited themselves to what was available from the previous study and have not engaged in any further chemical analysis. To be fair, this weakness is acknowledged (L462) and the need for further representative training data expressed as the next step. Additionally, given the quite detailed statistics, which use a Bayesian approach, some of the methodology might be made clearer to the more general reader as several concepts are not explained and taken for granted. Also, in places, the English is a little unclear or awkwardly expressed; I have tried to indicate some of the more difficult passages below. Finally, there are quite a few errors in the supplementary information.”

A1: We thank you for your critical appraisal of the scope of our manuscript. It is completely correct that it would have been nice if we would provide already in this study improved models which use peat material as training samples. As you point out, we do not hide this fact throughout the manuscript and fully acknowledge that this issue was discussed by a reviewer of the study of Hodgkins et al. (2018) (old: l. 53 to 54).

There are three pragmatic reasons and one scientific reason why the scope of the manuscript is limited to what we have analyzed and presented and we take the chance to describe these in more detail here, but also to justify the current scope of the study:

1. This manuscript is a by-product of a PhD project which was conducted with the aim to assess if the prediction models can be reliably used during this project. For this reason, there are limited resources to conduct additional measurements.
2. The required holocellulose and Klason lignin measurements would have been outside our area of expertise.
3. Considering the growing number of studies which use the original models of Hodgkins et al. (2018), we considered it important to publish our analysis as soon as possible to avoid future misinterpretation of the models’ predictions. An alternative option would have been to find collaborating partners who conduct (or conduct ourselves) the required measurements which would have cost additional time, but we would of course be open to that.
4. We think — as you do — that the manuscript in its present form makes genuine contributions to the scientific discourse, especially of researchers computing spectral prediction models for peat.

We build on the critical comment of a reviewer of the original study of Hodgkins et al. (2018) by giving these critical comments a quantitative, causal, and testable foundation.

The quantitative foundation is that our critical appraisal of the original models builds on an actual analysis of the models with data and thus provides more than what critical comments in a peer review process can and should reasonably do.

The causal foundation is that we provide — based on our quantitative analysis — a detailed causal explanation how and why the original training data are not representative for peat. It is easy to state a reasonable assumption that training data acquired from samples of a different type of material *may* not be representative for peat. We argue that it is a genuine contribution to show that this is actually the case and to explain why. From this analysis, we learn something about the chemistry of the original training data, the peat samples, and about the relation between this chemistry and the performance of prediction models (and spectral data). We think that this knowledge will be important also when peat samples are available as training samples, because there is no standard stating what is a representative dataset or stating how to define a set of sources as wide as necessary to compute models with a certain predictive accuracy. Peat itself can be quite diverse. We show here — as a hypothesis to be tested — what aspects probably are especially important when computing and validating improved models.

The testable foundation is that we provide hypotheses which can be tested. Our first hypothesis is that the normalized “carbohydrate” peak height (**carb**) is a good indicator for holocellulose content only when there are no (large) differences in OH bond types and abundances between the samples. We assume that relevant differences in OH bond types and abundances can also exist for peat samples because the chemistry and structure of cell walls of the peat forming vegetation can vary (both the amount and type of OH bonds and the holocellulose content) (e.g. comparing tropical peat with large fractions of wood, peat from (temperate) forest mires, or peat formed by *Sphagnum*). This hypothesis can be tested e.g. by comparing predictive performances of models fitted to either type of data. Testing this hypothesis would go beyond computing improved models: we can for example also learn something about how to interpret MIRS in terms of the decomposability of peat.

Our second hypothesis is that the normalized “aromatic” peak height (**arom15arom16**) is a good indicator for Klason lignin only when samples do not differ in protein content. In direct analogy to the “carbohydrate” peak height, this hypothesis is relevant (because also peat samples can differ in their protein content), it can be tested in the same way as our hypothesis for the normalized “carbohydrate” peak height, and testing it can make us learn something about peat chemistry and MIRS interpretation: we can for example learn whether **arom15arom16** (or intensities of similar parts of MIRS as are widely used in the interpretation or computation of humification indices (e.g. Broder et al. 2012)) are good indicators for some concept of decomposability. This may ultimately lead to a similar discussion as whether the C/N ratio is a good indicator for decomposability (or degree of decomposition), a discussion which to our knowledge has not been considered relevant yet.

In addition, we apply tools to validate models (particularly the analysis of the prediction domain, the analysis of model coefficients, the analysis of model residuals) which have to our knowledge only seldomly been applied on spectral prediction models for peat. In particular where sample sizes are small, such analyses could in analogy to ours here show how small peat datasets are also not representative. We hope that our study helps to circulate such model validation steps by showing in detail how they can be useful and how to apply them.

To summarize: We agree that it would be ideal if we had already today models which use peat training data to predict peat holocellulose and Klason lignin contents. However, we argue that our study makes genuine contributions which go beyond this aim and which merit publication.

We would also like to thank you for the detailed recommendations to improve language and method descriptions. We addressed these points in the specific comments below.

Specific Comments

Q2: “Abstract (L1) I feel there is a need to distinguish peat from soil organic matter. True, peat is a type of soil (a Histosol) and the organic matter of peat is a type of ‘soil organic matter’. However, the general expression of ‘soil organic matter’ is usually associated with mineral soils. Thus, in L3, Hodgkins et al. set

out to predict peat holocellulose and Klason lignin contents, not organic matter holocellulose and Klason lignin contents; indeed, they specifically removed the training samples containing silicates.”

A2: We assume that you suggest to include the word “peat” (old: l. 3 to 4; new: l. 3 to 4): “Recently, Hodgkins et al. (2018) developed models to quantitatively predict *peat* holocellulose and Klason lignin contents”. We think that this is a good suggestion to avoid confusion and we changed the line as suggested. There are several reasons why we framed the manuscript in terms of SOM and not exclusively peat and we list them here to explain our rationale:

1. We suggest that our study builds a bridge to studies computing prediction models for mineral soils because we show as a proof-of-concept that it is possible to predict holocellulose content also for the mineral-rich training samples. We therefore think that our study is also of interest for colleagues working with mineral soils.
2. There are studies computing prediction models which predominantly include mineral soil data, but nevertheless also include peat samples or even focus on peat samples (e.g. Helfenstein et al. 2021). These studies refer e.g. to soil organic carbon (SOC) which we also see as a term mostly used in the mineral soil community. However, as this study shows, these distinctions in terminology are less and less useful when soil C content is considered as continuum. We suggest that our study is of relevance also for mineral soils (because of our proof-of-concept mentioned above, because of the exemplary application of model validation approaches, and because of factors which may complicate prediction of holocellulose and Klason lignin contents which may also be relevant in mineral soils).

Q3: “L13 “mineral-rich samples” – it is unclear whether this refers to peat or to mineral soil. Of course, it could be both. As mentioned later (L96), some peats can become contaminated with mineral material to varying levels, though these are relatively rare occurrences.”

A3: We assume that you suggest to clarify already in l. 13 (old) what “mineral-rich samples” should refer to — mineral soil or peat. In fact, we wanted to communicate that this is a proof-of-concept applying to both cases. One problem is that we have no information on the mineral content for the particular training samples and another problem is that the training data probably are not representative for either case. However, the mechanism by which minerals create peaks in the spectra should be the same for either soil type.

We hope that the following change makes more clear that we refer to both soil types (old: l. 12 to 13; new: l. 12 to 14): “Finally, we provide a proof-of-concept that holocellulose contents can also be predicted for mineral-rich samples (e.g. peat with mineral admixtures or potentially mineral soils).”

Q4: “L29 “OM derived from plants and SOM” – better to reword as “SOM and OM derived from plants” (OM derived from SOM doesn’t make sense).”

A4: Thank you for pointing out this inconsistency. We actually wanted to write: “Plant and soil OM” (i.e. that the fractionation procedure is applied to both types of OM, not only SOM and not only plant OM). We changed the text accordingly (old: l. 29; new: l. 30).

Q5: “L37 “few species” – unclear if this is of wood or some other plants.”

A5: Thank you for pointing out that this is not clearly written. We changed the text to make this clearer (old: l. 37; new: l. 38): “Most of these models consider only wood, material from few woody and non-woody species, or only specific vegetation organs (Elle et al. 2019).”

Q6: “L46 “a later study” – need to specify which.”

A6: We here refer to all studies which were cited in the previous sentence. These are the studies which have cited Hodgkins et al. (2018) and which have used the models developed in Hodgkins et al. (2018) (and thus are the studies we are aware of which could have potentially validated the original models).

We changed the sentence to state this explicitly and hope that this was what you had in mind (old: l. 46; new: l. 48): “A major problem is that neither Hodgkins et al. (2018), nor any later study which was cited above and which used the models provided a thorough validation of the models.”

Q7: “L48 Replace “are” with “were”.”

A7: We are not entirely sure whether you refer to l. 48 or actually wanted to refer to l. 49. In l. 48 (old), the sentence “The peaks are baseline corrected, their maximum height computed and divided by the sum of absorbance values in the spectra.” refers to the procedure in general, and thus present tense should be correct. However, it is true that the next sentence (old: l. 49) “These values are used to calibrate the prediction models.” refers to what was done in Hodgkins et al. (2018) and thus past tense should be used (as mentioned in Q8).

We therefore changed the two sentences to (old: l.48 to 49; new: l. 50 to 52): “In the procedure, the peaks are baseline corrected, their maximum height computed and divided by the sum of absorbance values in the spectra. In Hodgkins et al. (2018), these values were used to calibrate the prediction models.”

Q8: “L49 Replace “are” with “were”.”

A8: Please see our reply to Q7.

Q9: “L51 Replace “for example” with “, for example,”.”

A9: Thank you. We changed the text as suggested (new: l. 54).

Q10: “L56 Replace “of” with “in”.”

A10: Thank you. We changed the text as suggested (new: l. 59). We made analogous changes in: old: l. 110, new: l. 117.

Q11: “L66 Replace “which” with “in which”.”

A11: Thank you for mentioning this issue. We think that “for which” is more appropriate here since we refer to assumptions of a distribution and not *in* a distribution (please correct us if you see this differently). We changed the sentence to (old: l. 66; new: l. 68): “In contrast, if unrealistic predictions occur, it is more reasonable to use a distribution for which assumptions are consistent with knowledge on the data generating process.”

Q12: “L67 I suspect most readers of ‘SOIL’ will not be so familiar with the term “beta distribution”; some explanation here is warranted.”

A12: Thank you for pointing this out. We provide a brief explanation in the subsequent sentence (old: l. 67): “For example, the beta distribution assumes a lower and upper bound which can be mapped to the interval [0, 100] mass-%.” This sentence mentions the values the Beta distribution is valid for and this is in our opinion the most relevant characteristic here.

To make clearer what the role of the beta distribution in a model is, we elaborated this to (correcting the interval the beta distribution is valid for) (old: l. 67; new: l. 70): “For example, the beta distribution assumes a lower and upper bound which can be mapped to the interval (0, 100) mass-%. It can be used as replacement for the normal sampling distribution used in ordinary linear regression models, such as the original models.”

We hope that this is an explanation you had in mind, but we are open for detailed suggestions on how to improve the text here.

Q13: “L77 The meaning of “prediction domain” should be given.”

A13: The term “prediction domain” is described in more detail in section 2.3. However, we agree that it is a good idea to explain the term already here: In principle, the prediction domain is the range of predictor variable values in the training data set.

To make this clear, we changed the sentences in l. 77 to 78 (old) to (new: l. 81 to 83): “*Is the prediction domain of the training data representative for peat samples?* If a model extrapolates outside the prediction domain — the range of predictor variable values in the training data set (Wadoux et al. 2021) — it is unclear if the predictions are valid.”

Q14: “L79 Replace “on” with “to”, replace “e.g.” with “, for example,”.”

A14: Thank you for pointing this out. We changed the text as suggested (new: l. 84). We made analogous changes in: (1) old: l. 197, new: l. 214; (2) old: l. 423, new: l. 446; (3) old: l. 426, new: 450.

Q15: “L80 The meaning of “no problem-specific correlations” is unclear.”

A15: Thank you. We agree that this term is not clear enough. With “problem-specific correlations”, we wanted to describe the dependence of a prediction model on spectral properties obviously specific for one type of material, only.

We changed the sentence to make clearer what we mean (new: l. 84): “An advantage of the independent reference sample set used by Hodgkins et al. (2018) is that the model is applicable to various OM types and, for example, does not depend on obviously material-specific spectral characteristics (Hodgkins et al. 2018).”

Q16: “L82 Generally, it is not good to start a sentence with “But”; “However,” might be better here.”

A16: Thank you. We changed the sentence as suggested (new: l. 87).

Q17: “L83 The “generality” of the training data set is not so obvious – it is quite wood based, since paper itself is derived from wood. Secondly, it could be argued to be a disadvantage since it overlooks any potential diagnostic peculiarities found only in peat.”

A17: We partly agree here. We agree that we should stress more that the generality is a *potential* advantage since it has not been tested so far. We tried to stress this aspect more as suggested (new: l. 88): “A potential advantage is therefore the potential generality of the training data set Hodgkins et al. (2018) used.”

However, we argue that the training data set is (at least one of the) most general training data sets available (based on the references cited in the introduction), and that the non-applicability to peat samples is also not so obvious without a detailed analysis.

It is true that the data set contains many wood samples, but it also contains not a small amount of grass, leave, and needle samples (22 % of the samples) which could potentially have similar spectral properties as peat samples. For this reason, we argue that without knowing our analysis, there are not many reliable arguments neither for, nor against the applicability of a model fitted to the data to peat samples.

Of course we agree with you that the original study should have validated the applicability of the training data, but we do not think that without any analysis it is possible to say beforehand that the models could not work. This is the only point we want to make here.

Q18: “L97 Replace “predictions” with “prediction”?”

A18: Thank you for pointing this error out. We assume that this refers to l. 87 and corrected this as suggested (new: l. 91).

Q19: “L98 Replace “also for” with “that includes””

A19: Thank you. We changed the text as suggested (new: l. 103).

Q20: “L99 Of course, as mentioned above, it would have been better to have a set of genuine mineral soil samples.”

A20: We agree that this important limitation should be stressed again here. We therefore added the following sentence (new: l. 105): “This is no replacement for a model using more training data, but it is a test whether it is at least in principle possible for MIRS to contain sufficient information to predict holocellulose contents of relatively diverse samples.”

However, we stress that it is not clear what a “set of genuine mineral soil samples” is unless one defines it via its prediction domain (or the molecular structures which eventually cause the prediction domain). This is probably not that easy since different minerals also differ in their spectral properties (e.g. Stuart 2004), meaning that this requires detailed research. Our study exemplarily shows approaches to analyze such difficulties.

Q21: “L114 “The data are available...” – delete this; it is obvious.”

A21: We agree that the formulation is not clear. We wanted to write that the data are accessible directly from the published paper (as supplementary information). It certainly is obvious that the data are used in Hodgkins et al. (2018), but we think that it is important to point to the exact data source used. Unfortunately, the supporting information has no own DOI.

To stress that we refer here to the access to the supplementary information, we changed the sentence to (new: l. 121): “The data are accessible from the supplementary information of Hodgkins et al. (2018).”

Q22: “L117 Replace “SOM” with “peat OM”.”

A22: We changed the text as suggested (new: l. 125).

Q23: “L122 “informative priors” – again this term may not be readily understood and should be made clearer. Secondly, what exactly these priors are should be made more specific.”

A23: On the one hand, we agree that it would be good to communicate this to readers unfamiliar with Bayesian analysis. On the other hand, we have to assume a certain familiarity with Bayesian statistics to avoid repeating too much content from textbooks on Bayesian statistics.

The term “informative” is no fixed term, but is a qualitative and problem-specific descriptor for the information content of a prior (Gelman, Simpson, and Betancourt 2017). Here, we mean with a “weakly informative prior” a non-flat proper (meaning that the prior probability distribution integrates to 1) prior which has no strong effect on the predictions in comparison to the respective frequentist model. That this was indeed the case was tested in supplementary information 1.

In particular, we assumed for both the intercept and regression coefficient a normal(0, 2.5) prior for z -transformed independent and dependent variable. This means that a one σ (denoting the standard deviation of z -transformed variables) change in the independent variable (e.g. the z -transformed normalized `carb` peak height) can result in the same change in the dependent variable (e.g. the z -transformed holocellulose content) when the absolute coefficient is ≥ 1 , which has a prior probability of approximately 0.69. The relevant information is that in comparison to the frequentist model, the Bayesian models produce the same prediction intervals.

To describe this better in the text, we changed the text to (old: l. 121 to 123, new: l. 128): “For this reason, to facilitate model comparison, we also recomputed the original models as Bayesian models with weakly informative priors. With weakly informative priors, we mean here priors which result in approximately the same prediction intervals as the original frequentist models (supplementary Fig. 1).”

We think that giving detailed parameter values would not support understanding unless you have more detailed knowledge on Bayesian statistics (and perhaps the software implementation) and we stress that these are available in full detail already via the reproducible research compendium (file `analysis/paper/002-paper-m-original-models.Rmd`, l. 40 to 41).

Q24: “L131 Replace “MCMC” with “Markov Chain Monte Carlo (MCMC)”.”

A24: Thank you. We changed the text as suggested.

Q25: “L139 Is there not a third way? Surely (holocellulose + Klason lignin) = 100%? In fact, it may be < 100% if there are other extracted components not included (I am not sure if this might be the case, since I do not have ready access to the De la Cruz et al. (2016) paper, describing the actual chemical analysis procedure). It is not clear if this constraint has been included within the statistical analysis.”

A25: Yes, you are right that (holocellulose + Klason lignin) $\leq$ 100 mass-% should hold, or more specifically, (holocellulose + Klason lignin + all other compounds) = 1 should hold. One could consider this with a Dirichlet regression model which is the multivariate extension to the beta regression model (e.g. Douma and Weedon 2019). We did not do this here for several reasons:

1. The analysis would become conceptually more complicated than it is now (which would make the analysis harder to understand) and we could not compare models for holocellulose and Klason lignin independently. This would also apply to the analysis of model coefficients.
2. The analysis would become computationally more difficult.
3. We are not aware of non-Bayesian approaches to Dirichlet regression for high dimensional data (e.g. PLSR) which would make such computations more feasible.
4. Not for all samples in the training data does (holocellulose + Klason lignin) $\leq$ 100 mass-% hold (see supplementary Fig. 4 (a), this is the case for sample type “office paper”). This implies that the data generating process is not fully compatible with the Dirichlet distribution (at least for contents as

measured for the training data). We assume that this implies that the measurement process needs to be improved since, apparently, the sum of the measured values cannot be unambiguously translated to mass-% in all cases. Note that this might also be a reason why Klason lignin contents for sample type “office paper” could not be appropriately modeled (Hodgkins et al. 2018).

However, we fully agree that this would be a useful improvement of our approach. To at least check whether the condition is violated for the peat and vegetation samples for the best binned beta regression models, we added individual MCMC samples from these models for holocellulose and Klason lignin. The histograms of the upper 95% prediction interval limits are shown in Fig. 1 below. The maximum value across all samples is 100 mass-%.

Both this maximum value and the histogram suggest that the new condition is not directly violated, but also that narrower prediction intervals can be expected from a Dirichlet regression model, since it is unlikely that holocellulose and Klason lignin contents are that high (i.e. 100 mass-%). Unfortunately, we are not aware of data for peat which would enable us to assess whether the median of the upper 95% prediction interval limit across all samples (73 mass-%) (or the median of the median, 58 mass-%) is sensible or not.

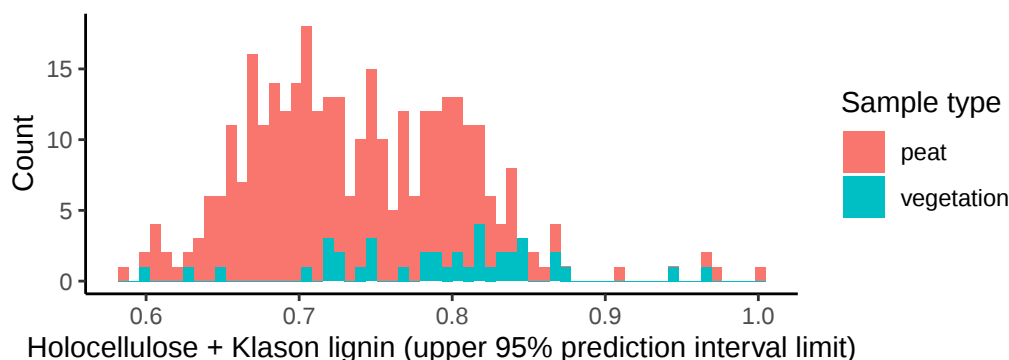


Figure 1: Histogram of the 95% prediction intervals for the sum of holocellulose and Klason lignin mass contents predicted with the best binned models.

To make clear that this constraint has not been included within the statistical analysis, that this is a limitation, and why the constraint was not considered, we added the following sentences to the text (old: l. 449, new: l. 474):

“A further limitation is that as the original models, the modified beta regression models do not consider the constraint that the contents of holocellulose, Klason lignin, and any remaining compounds should sum to 100 mass-%. This also represents a further test how realistic model predictions are (compare with Sect. 2.1). In principle, this constraint could be considered by using a Dirichlet regression model (e.g. Douma and Weedon 2019). We have not used this approach here due to the higher computational costs, potential computational difficulties, to keep the present model validation straightforward, and since the training data do not fulfill this condition for all samples (indicating that the measurement procedure needs to be improved, too).”

We did not add these sentences in l. 140 since in our opinion, this would have made the description of the methods unnecessarily complex.

Q26: “L140 Replace “in” with “within”.”

A26: Thank you. We changed the text as suggested (new: l. 148).

Q27: “L142 Replace “covering” with “covered”.”

A27: Thank you. We changed the text as suggested (new: l. 151).

Q28: “LL154-155 “in the form of depth profiles” – this is unclear; are you referring to actual peat depths here? Perhaps needs rewording.”

A28: Yes, we refer to plots where predicted values for peat samples are plotted against the depths of the peat samples in the peat cores. We were not sure what rewording you would suggest in this case.

We changed the text to hopefully make clearer what we mean (new: l. 163): “To facilitate practical comparison between the modeling approaches, we compared the predictions of all models for the peat samples by plotting predicted values versus depths of the peat samples.”

Q29: “L165 Replace “amont” with “amounts”. The term “regularization” may need some defining. Replace “for” with “of”?”

A29: Thank you for pointing out these errors. We changed the text as suggested (new: l. 174).

Again, we had to face a trade-off between writing a short and concise manuscript which focuses on what we did and explaining in more detail statistical terms which may not be known to all readers. We hope that the following additions make the text better understandable (and we would like to thank you for mentioning such aspects!) (new: l. 174):

“To avoid overfitting, we computed models using priors implying different amounts of regularization — shrinking coefficients to 0 with the aim to avoid overfitting — of model coefficients (standard deviation of 1 and 0.5, respectively) ...”

Q30: “L168 It might have been interesting to compare the approach given in 2. to a PLS (Partial Least Squares) analysis, which has been applied to similar data (see Artz et al.) and is a standard (non-Bayesian) approach for such multivariate data.”

A30: It is true that there are many more modeling approaches which probably would have a similar predictive accuracy as the approach we used here (especially since the sample size is small). For example, reviewer 1 suggested principal component regression (PCR) as yet another alternative approach and there are many more approaches one could have tried (e.g. supervised PCR, iterative supervised PCR, interval PLSR, moving window interval PLS, etc; see e.g. Xiaobo et al. (2010)).

Based on a previous study (Teickner, Gao, and Knorr 2022), we assume that for datasets with a similarly small sample size (around 50), the regression approach we used here and partial least squares regression (PLSR) often have similar predictive accuracies. Of course, this depends on the particular dataset used. Based on this previous study, we assume also here that PLSR would have improved the predictive accuracy in a similar way as our “best binned” models and would potentially have reduced the (potential) biases we found for the original models.

One advantage of the approach we used in comparison to dimension reduction approaches (e.g. PCR, PLSR), which summarize the original predictors into latent variables, is that coefficients for individual predictors are estimated more independently, since no latent variables are computed as summary of all predictors. Therefore, multiple regression with regularizing priors has the advantage of facilitating model interpretation (see e.g. Fig. 3).

A practical disadvantage we would have faced when using PLSR is that we are not aware of a fully working R implementation of a beta PLSR approach. This would have made, in our opinion, the implementation of some of the improvements we suggest here using PLSR difficult and/or conceptually similarly difficult as the Bayesian approach we used here.

We do not claim that our approach is perfect yet, nor that we have filled all research gaps (however, we certainly have addressed the research questions set out in the introduction). We welcome any studies which want to further improve our models or compare them to other modeling approaches (all models can be reproduced via the reproducible research compendium).

We think that it is good to at least briefly mention these aspects to provide readers a better orientation. To this end, we added the following paragraph (old: l. 175; new: l. 185): “An alternative, popular, approach to approach 2 would be dimension reduction, for example via partial least squares regression, principal component regression, or variants of these (Xiaobo et al. 2010). In general, there are many alternative approaches which could be tested to use more information contained within the spectra than the original models, and many of these probably would result in similar predictive performances as approach 2 (regularization), especially when sample sizes are small (Xiaobo et al. 2010; Teickner, Gao, and Knorr 2022). An advantage of regularization is that model coefficients are estimated more independently than in dimension reduction approaches, which makes it more straightforward to interpret model coefficients. The key is that the approaches we chose are suitable to analyze our research questions.”

Q31: “Figure 2. The key would be better expressed as “Bayesian models $\hat{\alpha}$; Beta $\hat{\alpha}$; Gaussian”, since both

are Bayesian. The sentence “Dashed grey...” could be omitted; it is pretty obvious.”

A31: Thank you. We changed the figure as suggested (new: Fig. 2).

Q32: “L276 Define the “fingerprint” region (it could be shown in Figure 1).”

A32: We changed the sentence to (new: l. 293): “Instead, bins with large absolute coefficients are located in the fingerprint region (600 to 1500 cm^{-1}) and probably related to aromatic in plane C-H bending (~ 1150 and $\sim 1270 \text{ cm}^{-1}$) (Stuart 2004) ...”.

We think that it is not necessary to visualize this range in a figure as the term was just used to indicate the rough location of the selected bins. The fingerprint region is defined based on Stuart (2004).

Q33: “Figure 3. Under “Holocellulose (no minerals)”, the blue line at ca. 1590 (negative coefficient > 0.2) is unmarked – should it be?”

A33: Yes, it should be unmarked because the median absolute coefficient is not > 0.2 (it is slightly smaller).

Q34: “Table 1. The column “Original model?” is not helpful, more confusing – omit (given in legend, whether .2 or .3).”

A34: Thank you for this suggestion. We changed the table accordingly (new: Tab. 3).

Q35: “LL304-305 This is really unsurprising.”

A35: You are right that one *can* have a reasonably strong assumption/belief that it is *a priori* unsurprising that the prediction domain of the training dataset does not cover the peat samples. We also explicitly described this assumption in the introduction (old: l. 84 to 86; new: l. 89 to 91) showing that we were not too surprised, too.

However, in our opinion, a mere assumption needs to be backed up by a quantitative analysis in order to be suitable as an argument to suggest concrete improvements. We cannot criticize existing models by reasonable assumptions, we need tests of these assumptions.

Q36: “LL322-323 Again, this points to the training dataset not being suitable for the task in hand.”

A36: Yes, this is correct (as we mention in the manuscript (old: l. 327 to 329, 356 to 357; new: l. 349 to 351, 378 to 379)).

However, we stress that our analysis does more than this because it explains why the training data are not suitable, i.e. which characteristics of the training samples cause which spectral features which cause the bias. This knowledge is important to choose suitable training data and to check under which conditions models computed with such data do make accurate predictions.

We do not think that collecting a “genuine” training dataset is a trivial problem, exactly for the key arguments we provide in our study. These key arguments are that we need to pay attention to differences in protein content and to differences in other major spectral features, even if we do not include them as predictors into the model.

Q37: “L360 “classes” – this is a new term which has not been defined. We are left to guess what these classes are. They are given in Supplementary Figure 11, but at this point we are yet to read this.”

A37: After answering comment Q58 below, we noticed that the supplementary text and also the main text are written ambiguously, since we do not list all sample types belonging to class 1.

To address this valuable point, we changed the main text as follows (old: l. 383 to 384; new: l. 405 to 406): “This is what distinguishes samples of class 1 and 2 in the training data (supplementary Fig. 11): Samples of class 1 are wood samples and paper product samples derived from wood.”

We think that this resulted in a misunderstanding. The classes are defined in the sentences following l. 360 (old): “Samples of these classes differ in the relative contribution of the **arom15** and **arom16** peaks to **arom15arom16** (the sum of their heights) (Fig. 6). The first class (class 1) has high contributions of the **arom15** peak ($> 30\%$) and typically smaller **arom15arom16** values; the second class (class 2) a smaller contribution of the **arom15** peak and larger **arom15arom16** values.”

In the main text, we refer to Fig. S11 in the context of the two classes only in l. 383 (old) (new: l. 405) and there, we describe that the two classes comprise different sample types. This order is due to how we defined

the classes and how we build up our argument: first we describe what spectral features are related to the bias and then we explain how the bias arises with the help of knowledge of additional sample properties. Thus, this order is by intention.

We hope that this makes more clear how the classes are defined, that we provide these definitions in the main text, and sample types are not the same as the two sample classes. We apologize that this has been written confusingly. If this is still confusing, we would be glad to receive suggestions for how to make this clearer.

Q38: “L395 “because they interfere less with” – no, the reverse applies: “because they are interfered less by.””

A38: We see interference as a symmetrical interaction. But we see that in this context the wording could be improved as suggested. We changed the text as suggested (new: l. 418). Thank you.

Q39: “L396 Replace “that” with “why.””

A39: Thank you. We changed the text as suggested (new: l. 419).

Q40: “Figure 6 “If the points are...” – I think you can omit this; it is a rather trivial remark and obvious to all.”

A40: We changed the figure caption as suggested (new: Fig. 6).

Q41: “L403 The ELPD values quoted seem to be the same – is this correct?”

A41: Yes, this is correct, but it is true that we could have written this less confusingly. The two statements say the same thing (the first says that the model using binned spectra had an average ELPD 17.24 larger than the model using all peaks; the second says that ΔELPD — the difference in average ELPD of the model using all peaks relative to the model using binned spectra is -17.24).

We changed the text by deleting the first statement (new: l. 425). Thank you for pointing this out.

Q42: “L422 Replace “to predict” with “being predicted.””

A42: Thank you. We changed the text as suggested (new: l. 445).

Q43: “L423 Replace “on” with “to.””

A43: Thank you. We changed the text as suggested (new: l. 446).

Q44: “L427 Replace “on than” with “to than.””

A44: Thank you. We changed the text as suggested (new: l. 450).

Q45: “L429 Replace “what” with “as to what.””

A45: Thank you. We changed the text as suggested (new: l. 452).

Q46: “L439 Replace “problem is that also” with “problems are that.””

A46: Thank you for pointing out this error. We changed the text (new: l. 462) to: “The most important problem is ...” (as the sentence refers to one problem).

Q47: “L442 Replace “probably ca” with “can probably”. However, I think “probably” is inaccurate here; I think you say “certainly.””

A47: Thank you for pointing out this error. You are also right that we mean “certainly”. We changed the text as suggested (new: l. 465).

Q48: “L444 “calibration transfer” needs definition or more explanation.”

A48: Thank you. We changed the sentence to make this clearer (new: l. 465): “(2) It is unclear how robust the original models and our improved models — especially models using binned spectra — are in terms of calibration transfer (calibration transfer is the application of a model to spectra measured differently than the training data, e.g. with a different procedure, on a different device, in a different laboratory (Workman 2018)).”

Q49: “L450 I suggest replacing “SOM” with “actual peat samples” – the idea of application to SOM (in general and including mineral soils) comes two sentences later. Hence in L451 I suggest omitting “SOM and”.”

A49: Thank you. We changed the sentence to (new: l. 481): “In summary, our analysis opens concrete and promising directions to further improve the models: We need training and validation data that include peat, particularly highly decomposed peat.”

Q50: “L464 “interpreted with caution” – it would be useful to develop this thought more. The estimates given in Hodgkins et al. have been applied to give far-reaching conclusions regarding peat carbon storage, latitudinal trends and possible responses to climate change. Using the same basis, ideas have been amplified and extended in the Verbeke et al. (2022) paper. Can we still rely on the overall findings from these papers or are the conclusions now in doubt? I suspect that the revised models given here, while improving the accuracy of determinations to some minor degree, will not change the overall pictures presented in these (and perhaps other) derivative papers but some discussion on this would be very helpful and would give a wider context and application.”

A50: We appreciate this suggestion as we also think that further discussions on this proposed gradient are useful.

Indeed, we have elaborated this point further (and added additional remarks and suggestions regarding uncertainty handling) in a companion manuscript intended as a reply to Hodgkins et al. (2018) (this manuscript was written before Verbeke et al. (2022) was published). This reply manuscript is not yet published, since we first wanted to publish the manuscript discussed here which represents the technical basis for our reply manuscript. However, the script for the reply manuscript (containing also the text) is available via the reproducible research compendium to the manuscript discussed here.

In brief, the results differ if the modified models are used due to the strong underestimation of Klason lignin contents (i.e. Klason lignin contents are strongly underestimated in Hodgkins et al. (2018)). As a consequence, also the depth-related relative accumulation of Klason lignin during decomposition is probably underestimated. We think that this is an important finding since a larger amount of Klason lignin is required according to the modified models for the hypothesized stability of tropical peat and deeper peat than has been suggested.

Another point is that because of the bias in the models, the original results can only be interpreted in terms of qualitative concepts for carbohydrates and aromatics which makes it much more difficult to connect the results to other studies, to generalize them, and to test them.

To foreshadow more explicitly that we are preparing such an analysis, we modified the sentence (new: l. 496): “Results from currently published studies using the original models should be interpreted with caution (we are currently preparing a manuscript, see Teickner and Knorr (2022), which explores the implications of our results here for the results of Hodgkins et al. (2018)).”

Q51: “L466 Delete “, also in practice”.”

A51: We changed the text as suggested (new: l. 500).

Supplementary Information

Q52: “A general comment on the information here: apart from the introductory paragraphs to S1 and S2, there is a description of the figure followed by the legend to the figure and often the material is a repeat. Really, it would be better just to have a descriptive legend.”

A52: Thank you for making this suggestion. We changed the text as suggested.

Q53: “S1 Replace “as baseline” with “as a baseline”.”

A53: Thank you. We changed the text as suggested.

Q54: “Figure 3. See same comment as for Figure 2 above.”

A54: Thank you. We changed the figure as suggested (new: Fig. S3).

Q55: “Figure 4. The axis labels for (a) are the wrong way round! For “Sample type” is this the same as “class” in the main text? If so change to “class” (but see on Figure 11 below).”

A55: Thank you for pointing out this error. The axes labels are indeed the wrong way round. “Sample type” is not the same as “class” (the two classes defined in the manuscript text) in the main text. We changed the figure as suggested (new: Fig. S4).

Q56: “Figure 6. Replace “overestimation” with “underestimation” and “underestimation” with “overestimation”.”

A56: Thank you for pointing out this error. We changed the caption as suggested (new: Fig. S6).

Q57: “Figure 8. I would suggest having Training, Peat and Vegetation across the top and peak heights down the side (as laid out in Figure 7). Also (a) and (b) seem to show the same data; can this be right? In the legend it says peat (second row) but it is actually peat in the third row.”

A57: Thank you for spotting this error! Indeed, (a) did not show the data for the intensity at 3400 cm^{-1} , but for 1250 cm^{-1} . We corrected this and also changed the figure as suggested (new: Fig. S8).

Q58: “Figure 10. Replace “6.31e-16” with “0”.”

A58: Thank you! We changed the figure as suggested (new: Fig. S10).

Q58: “Figure 11. Replace “type 2” with “class 2”. It is unclear what these classes are, and where does paper and cardboard fit in? Perhaps indicate in the list of sample types which are class 1, 2, 3, etc.”

A58: Thank you. Since this text is deleted (see A52), we did not change the text. However, it is true that this was not described clearly enough and we think that this might also be relevant for the misunderstanding (we assume) for Q37. We hope that this comment was addressed in A37.

Q59: “Figure 12. The order of regressions mentioned in the description is not the same order as actually depicted. In the legend, replace “binnes” with “binned”, “50” with “20” (I presume), and “line” with “lines”.”

A59: Thank you for pointing out these errors. We changed the caption as suggested (new: Fig. S12).

Q60: “Figure 13. The “Dataset” key could be omitted – it is superfluous.”

A60: Thank you for the suggestion. It is true that the key is redundant. However, we would like to keep the redundancy here since the color scheme is applied across different figures and we think that this helps recognizing the different dataset types. If you think this is not only superfluous, but even confusing, we will remove the key here.

Additional changes

We made the following additional changes:

1. old: l. 4; new: l. 4: We changed the sentence “The models may have the potential to understand large-scale SOM gradients and have been used in various studies.” to “The models may help to understand large-scale SOM gradients and have been used in various studies.”
2. old: l. 40; new: l. 42: We changed “comprises” to “comprise”.
3. old: l. 47; new: l. 49: We replaced “the authors” by “Hodgkins et al. (2018)”.
4. old: l. 49; new: l. 51: We replaced “... and divided by ...” by “..., and normalized — divided by ...” to introduce the term “normalization” which is used later in the text and also in Hodgkins et al. (2018).
5. old: l. 52; new: l. 55: We replaced “data is” by “data are”.
6. old: l. 65 to 66; new: l. 68: We removed “due the central limit theorem” because this applies only to mean values. However, we would see the application of a normal distribution only as practically justified when predictions for individual samples would not generate unrealistic values and prediction intervals.

7. old: l. 91; new: l. 96: We changed “models computed” to “computed models”.
8. old: l. 151; new: l. 160: We replaced “data is” by “data are”.
9. old: l. 168; new: l. 178: We changed “(Teickner, Gao, and Knorr 2022)” to a textual citation.
10. old: l. 203; new: l. 220: We changed “represent” to “form”.
11. old: l. 218; new: l. 235: We changed “We therefore created such plots” to “We therefore plotted measured values against predicted values”.
12. old: l. 294; new: l. 316: We changed “samples — and” to “samples and —”.
13. old: l. 307; new: l. 329: We replaced “data is” by “data are”.
14. old: l. 328; new: l. 350: We replaced “data is” by “data are”.
15. old: l. 352; new: l. 374: We removed “middle row in” as the described pattern is visible in rows 2 to 4 of the figure.
16. old: l. 414; new: l. 437: We changed “the holocellulose content” to “holocellulose contents”.
17. old: l. 423 to 425; new: l. 444 to 448: We replaced “Most importantly, due to spectral normalization, predictions can even be sensitive to variables not included in the model. Therefore, to assess if training data (the prediction domain) is representative, whole spectra have to be compared.” by “Most importantly, due to spectral normalization, predictions can be sensitive *even* to variables not included in the model. Therefore, to assess if training data (the prediction domain) *are* representative, whole spectra have to be compared.” (changes emphasized). Additionally, we replaced “data is” by “data are” in the previous sentence.
18. old: l. 430; new: l. 453: We replaced “for” by “to”.
19. old: l. 444; new: l. 468: We changed “We assume the binning and estimation procedure is relatively robust ...” to “We assume that the binning and estimation procedure are relatively robust ...”
20. old: l. 446; new: l. 471: We changed “if this holds in general” to “whether this is possible in general”.
21. old: l. 456 to 457; new: l. 488 to 490: We corrected this statement to “To support such developments, we implemented the best models using binned spectra into the R package *irpeat* (Teickner and Hodgkins 2020). The other models can be reproduced from the reproducible research compendium (Teickner and Knorr 2022).”
22. old: Tab. 1; new: Tab. 1: In the caption, we changed “ Δ ELPD the difference in ELPD relative to the best average model” to “ Δ ELPD the difference in ELPD relative to the average ELPD of the on average best model”.
23. old : Fig. S2; new: Fig. S2: In the caption, we changed “all absorbance value” to “all absorbance values”.
24. old : Fig. S2; new: Fig. S2: We improved the caption to make clearer what the different panels show.
25. old : Fig. S2; new: Fig. S2: We improved the text above this figure make clearer what the implication of the simulation is.
26. We relabeled all supporting figures to the “S + number” convention (e.g. “supporting Fig. S3” instead of “supporting Fig. 3”).
27. We adopted the color scheme of figures where these were not recognizable for color blind people.

References

Broder, T., C. Blodau, H. Biester, and K. H. Knorr. 2012. “Peat Decomposition Records in Three Pristine Ombrotrophic Bogs in Southern Patagonia.” *Biogeosciences* 9 (4): 1479–91. <https://doi.org/10.5194/bg-9-1479-2012>

1479-2012.

De la Cruz, Florentino B., Jason Osborne, and Morton A. Barlaz. 2016. "Determination of Sources of Organic Matter in Solid Waste by Analysis of Phenolic Copper Oxide Oxidation Products of Lignin." *Journal of Environmental Engineering* 142 (2): 04015076. [https://doi.org/10.1061/\(ASCE\)EE.1943-7870.0001038](https://doi.org/10.1061/(ASCE)EE.1943-7870.0001038).

Douma, Jacob C., and James T. Weedon. 2019. "Analysing Continuous Proportions in Ecology and Evolution: A Practical Introduction to Beta and Dirichlet Regression." Edited by David Warton. *Methods in Ecology and Evolution* 10 (9): 1412–30. <https://doi.org/10.1111/2041-210X.13234>.

Elle, Oliver, Ronny Richter, Michael Vohland, and Alexandra Weigelt. 2019. "Fine Root Lignin Content Is Well Predictable with Near-Infrared Spectroscopy." *Scientific Reports* 9 (1): 6396. <https://doi.org/10.1038/s41598-019-42837-z>.

Gelman, Andrew, Daniel Simpson, and Michael Betancourt. 2017. "The Prior Can Often Only Be Understood in the Context of the Likelihood." *Entropy* 19 (10): 555. <https://doi.org/10.3390/e19100555>.

Helfenstein, Anatol, Philipp Baumann, Raphael Viscarra Rossel, Andreas Gubler, Stefan Oechlin, and Johan Six. 2021. "Quantifying Soil Carbon in Temperate Peatlands Using a Mid-IR Soil Spectral Library." *SOIL* 7 (1): 193–215. <https://doi.org/10.5194/soil-7-193-2021>.

Hodgkins, Suzanne B., Curtis J. Richardson, René Dommain, Hongjun Wang, Paul H. Glaser, Brittany Verbeke, B. Rose Winkler, et al. 2018. "Tropical Peatland Carbon Storage Linked to Global Latitudinal Trends in Peat Recalcitrance." *Nature Communications* 9 (1): 3640. <https://doi.org/10.1038/s41467-018-06050-2>.

Kubo, Satoshi, and John F. Kadla. 2005. "Hydrogen Bonding in Lignin: A Fourier Transform Infrared Model Compound Study." *Biomacromolecules* 6 (5): 2815–21. <https://doi.org/10.1021/bm050288q>.

Piironen, Juho, Markus Paasiniemi, and Aki Vehtari. 2020. "Projective Inference in High-Dimensional Problems: Prediction and Feature Selection." *Electronic Journal of Statistics* 14 (1). <https://doi.org/10.1214/20-EJS1711>.

Stuart, Barbara H. 2004. *Infrared Spectroscopy: Fundamentals and Applications*. Analytical Techniques in the Sciences. Chichester, UK: John Wiley & Sons, Ltd. <https://doi.org/10.1002/0470011149>.

Teickner, Henning, Chuanyu Gao, and Klaus-Holger Knorr. 2022. "Electrochemical Properties of Peat Particulate Organic Matter on a Global Scale: Relation to Peat Chemistry and Degree of Decomposition." *Global Biogeochemical Cycles*, February. <https://doi.org/10.1029/2021GB007160>.

Teickner, Henning, and Suzanne B. Hodgkins. 2020. "Irpeat: Simple Functions to Analyse Mid Infrared Spectra of Peat Samples."

Teickner, Henning, and Klaus-Holger Knorr. 2022. "Hklmirs: Reproducible Research Compendium for "Improving Models to Predict Holocellulose and Klason Lignin Contents for Peat Soil Organic Matter with Mid Infrared Spectra" and "Predicting Absolute Holocellulose and Klason Lignin Contents for Peat Remains Challenging"," March. <https://doi.org/10.5281/ZENODO.6325760>.

Verbeke, Brittany A., Louis J. Lamit, Erik A. Lilleskov, Suzanne B. Hodgkins, Nate Basiliko, Evan S. Kane, Roxane Andersen, et al. 2022. "Latitude, Elevation, and Mean Annual Temperature Predict Peat Organic Matter Chemistry at a Global Scale." *Global Biogeochemical Cycles*, January. <https://doi.org/10.1029/2021GB007057>.

Wadoux, Alexandre M. J.-C., Brendan Malone, Budiman Minasny, Mario Fajardo, and Alex B. McBratney. 2021. *Soil Spectral Inference with R: Analysing Digital Soil Spectra Using the R Programming Environment*. Progress in Soil Science. Cham: Springer International Publishing. <https://doi.org/10.1007/978-3-030-64896-1>.

Workman, Jerome J. 2018. "A Review of Calibration Transfer Practices and Instrument Differences in Spectroscopy." *Applied Spectroscopy* 72 (3): 340–65. <https://doi.org/10.1177/0003702817736064>.

Xiaobo, Zou, Zhao Jiewen, Malcolm J. W. Povey, Mel Holmes, and Mao Hanpin. 2010. "Variables Selection Methods in Near-Infrared Spectroscopy." *Analytica Chimica Acta* 667 (1-2): 14–32. <https://doi.org/10.1016/j.aca.2010.03.048>.