

The Many Dimensions of Uncertainty Analysis in LCA

By Gregory Norris

Uncertainty comes into play when analysts, databases, or reports make statements or assertions about something in the real world based on Life Cycle Assessment methods and data. These statements or assertions can be either quantitative or qualitative. An example of a quantitative statement within LCA might be, “the total ‘embodied’ carbon dioxide emissions caused by generating a kilo-watt hour in the northeastern US are x lbs per kWh.” An example of a qualitative assertion might be, “the life cycle particulate emissions of recycled newsprint are lower than those of newsprint made from virgin fibers.”

Why do we care about uncertainty in LCA? Because the assertions or statements we make about the world based on LCA may be wrong — they are uncertain. We presumably use the results of LCA to help us, or others, decide future courses of action. Future decisions based (at least in part) on our LCA analysis will be better decisions if they take into account not only our results and conclusions, but also the uncertainty in these results and conclusions. This point was made rather forcefully by Wilson et al [1985]: “A decision made without taking uncertainty into account is barely worth calling a decision.”

Useful information about uncertainty in LCA will tell us something about the likelihood or probability that our statements and conclusions are right or wrong. Ideally, it will also establish “confidence bounds” on our results — “wobble room” around our results within which the true values have an estimated likelihood of falling.

For example, consider our example LCA-based assertion that, “life cycle particulate emissions of recycled newsprint are lower than those of newsprint made from virgin fibers.” This statement may be right or wrong. How *reliable* is it? How *confident* are we that it indeed holds true? The terms **confidence** and **reliability** go hand in hand. That is, when we assert that our conclusion is 90% reliable, this is the same as asserting that we can have 90% confidence in our results.

In classical statistics, “90% confidence” means that if we were to repeat our analysis very many times, each time using new and equally probable, randomly selected values for our uncertain quantities, our conclusions would be correct 90% of the time. We cannot use *all* the methods of classical statistics in LCA, primarily because our underlying LCA data are not based on random samples, *and* we are not strictly dealing with “random variables” which follow some known or unknown frequency distribution. However, we *are* dealing with uncertain quantities, about which we may develop or obtain **subjective probability distributions**, so we *can* use methods from the well-developed and active fields that fall under the broad heading of **uncertainty analysis**.

Subjective probability distributions are an alternative to simple point estimates. Rather than use in our analysis a single number to estimate some real-world quantity (say, tons of sheet steel required to make one ton of product), we instead use a range within which we expect the true value to lie, and optionally a description of the relative likelihoods (probabilities) that the true value lies within certain portions of this range.

IF we use or develop subjective probability distributions in this manner, we can then use the methods of uncertainty analysis in LCA and truly take uncertainty into account in our results, conclusions, and decisions. *THEN*, when we make an LCA-based statement such as, “life cycle particulate emissions of recycled newsprint are lower than those of newsprint made from virgin fibers,” we will also be able to understand, communicate, and take into account the reliability of the statement. If the reliability of the conclusions is not sufficient for our decision-making needs, uncertainty analysis will also help us identify which data uncertainties are most influential. It can further help us determine the levels of reduction in data uncertainty required to reach a specified level of results reliability or confidence, and/or it can tell us what the results reliability consequences will be of pre-specified reductions in data uncertainty. In short, we will no longer be “flying blind” about uncertainty and reliability in LCA.

Uncertainty Analysis versus Data Quality Characterization in LCA

Most efforts directed at uncertainty in LCA to date have focused on a related issue: **data quality**. We review some of these efforts in a later section of the report. As generally used in LCA, the term “data quality” is a very broad heading which might be translated, “facets of goodness and other characteristics of the data.” Practitioners generally elaborate different dimensions of data quality, or **data quality indicators**, from five to eighteen or more in number. They suggest that individual **data elements** be evaluated, and possibly scored, with respect to each data quality dimension or indicator. In most cases [notably RTI 1992, and EPA 1995], authors then suggest that the quality of the data, as expressed in the data quality indicators, be compared with **data quality goals** in order to assess the adequacy of the data quality.

Data quality is a valid and potentially useful topic for consideration. However, there are some very important problems with taking a data quality-based approach to analyzing uncertainty in LCA data and results. The primary problem is in determining what the implications of data quality indicators are for the quality, and reliability, of LCA results and conclusions. The difficulty stems from two major factors: the need to synthesize across dimensions of data quality and across data elements, and the influence of context (usage) upon uncertainty.

1. The need to synthesize across dimensions of data quality and across data elements. When several different data quality dimensions or indicators are used, what is the net influence or implication of these separate facets of data quality upon the uncertainty, or reliability, of a given data element? Further, what are the implications of the quality of individual data elements for the reliability of results as a whole?

Many authors have usefully pointed out that there are separate dimensions of quality and reliability in data [Funtowicz and Ravetz 1987, 1990; Pate-Cornel 1996; Cullen and Frey 1999]. It can be very helpful to identify and preserve information about these distinctions. But we *also* need to understand the “bottom line” implications of data quality and uncertainty upon our data, our results, and especially our conclusions.

When only qualitative information about data quality is elicited from experts, we are highly suspect in trying to do something quantitative with it. Indeed, without carefully and empirically verified methods for translating qualitative assessments into their quantitative implications for uncertainty, we will be misusing probability distributions and uncertainty analysis in these instances, and we will be grossly in error if we try to make assertions about results reliability from such a basis. Much better to ask the experts to include explicit assessments of uncertainty up front.

Also, just knowing that half of our data is of good quality, and half is not, tells us nothing about the quality — let alone the reliability — of our results. For example, some data elements play highly influential roles in our analysis, while others are quite insignificant. If electricity generation is the major source of particulate emissions in the life cycle, then reliable data on this emission factor is more critical for this process than for other processes in the study.

2. The influence of context (usage) upon uncertainty. Uncertainty (reliability) is partly a property of the data itself, but largely it arises through *usage*, through application in a particular modeling exercise.

What is the uncertainty in our data for particulate emissions from coal-fired power plants? It depends largely upon *which* coal-fired power plants we have data for, *and which* coal-fired power plants we are using this data to model! The relationship between the subject of the data and the object of our modeling is critical. It can even be argued that if we are using the data for a given power plant to model *that* power plant, the uncertainty in our estimates still depends on the age of the data in relation to the time frame we are modeling.

Thus, we cannot assess once and for all the uncertainty or reliability of a given data element in LCA. The implications of that data’s usage for the uncertainty in our conclusions depends largely upon *how* we use it — that is, *what* we use it to model. We address this issue in a subsequent section. First, we investigate the types of data and results that are involved in LCA.

Categories of Result Types in LCA

We begin by noting that the “results” of one LCA may serve as an element in the database underlying another LCA. For example, cradle-to-gate Life Cycle Assessments of different plastic resins report cradle-to-gate summary life cycle inventories for these resins, per functional unit (e.g., per kilogram of resin in pellet form). These cradle-to-gate inventories often become single entries in an LCA database, used by other practitioners as among the basic “data building blocks” of their studies. The important points for our purposes are that —

- a) not all “LCA data” — that is, data in databases from which LCAs are created — are process-level data; and
- b) both process-level data and cradle-to-gate data are “assertions” or “hypotheses” within LCA, about which uncertainty naturally arises.

Thus, both are “results” and “data” at the same time. Both can be building blocks for subsequent LCAs. Both are assertions about physical reality, as seen through the methodological lens of LCA.

Six life cycle inventory (LCI)¹ results types can be defined using two dimensions (see Figure 1). First, we distinguish between different levels of *scope*:

- a) *process level* data, assertions, or results;
- b) *tree-level* data, assertions, or results; and
- c) *life cycle level* data, assertions, or results.

Next, for each of these levels of scope, results can be presented in two forms: either as statements about a single process, process tree, or life cycle; or as comparisons between two processes, trees, or life cycles.

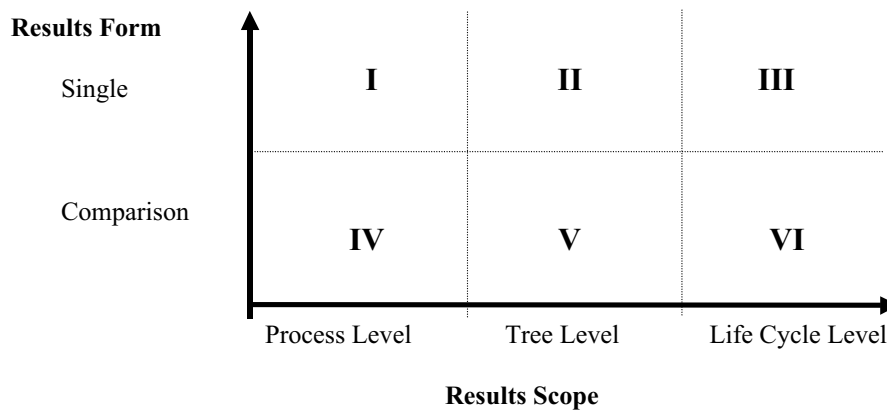


Figure 1: Six types of LCI results

Process Level Results: Examples of process level results include —

- quantities of each pollutant released (to the environment) per unit of process output;
- quantities of each raw material consumed (directly from the environment) per unit of process output;
- quantities of each intermediate material consumed (from other unit processes rather than directly from the environment) per unit of process output; and
- quantities of each co-product generated per unit of process output.

¹ For simplicity, we focus the remainder of the report on uncertainty in Life Cycle Inventory analysis. Consideration of uncertainty in LCI is a pre-requisite for subsequent efforts directed at dealing with uncertainty in Life Cycle Impact Assessment

Tree Level Results: The term “tree” refers to the whole chain of processes whose outputs are used either directly or indirectly by a given “driving process.” For example, if the driving process is electricity production, then the driving product is electricity, and its tree is the full set of processes (coal mining, petroleum extraction and refining, and transport, etc.) whose outputs are needed either directly by the generating plant, or by suppliers to it, or by suppliers to them, and so on. The tree is *quantitative* as well as structural, in that it reflects *how much* of each output is required from each process in the tree, all playing their parts to enable the production of one unit of driving product from the driving process.

Trees may be “upstream” (also called “cradle-to-gate”) trees, or “downstream” (also called “gate-to-grave”) trees. Downstream or gate-to-grave trees are the full set of processes which support the end-of-life processing routes relevant to the given product output (generally a combination of recycling, landfilling, and incineration, together with transportation and other processes supporting these end-of-life management processes).

What distinguishes tree-level results (or data) from process-level results (or data) is the fact that tree-level results are meant to be, and are taken to be, estimates of the *full* embodied or induced inventories associated with production of the driving product. For upstream or cradle-to-gate trees, these data are the result of attempting to trace all the way back through other upstream unit processes to the raw material extraction activities induced by this process tree or “process chain.” These tree-level modeling results are the outcome of life cycle inventory computations and are influenced in part by methodological assumptions made in the course of carrying out the computations (such as co-product allocation methods, boundary truncation conventions, etc.).

Life Cycle Level Results: Life cycle results are generally composed of sets of multiple upstream and downstream tree-level results. We will see that trees and life cycles are treated quite similarly with regard to uncertainty issues — and both quite differently from single processes.

Comparisons vs. Statements about a Single Process/Tree/Life Cycle: The processing and reporting of uncertainty in final results must be done differently depending upon whether the final result is a *statement* about a single process/tree life cycle, or a *comparison* between processes/trees/life cycles.

For example, it is possible to be highly uncertain about the life cycle carbon dioxide emissions of virgin paper, and of recycled paper, yet at the same time to be highly *certain* about which has the higher life cycle carbon dioxide emissions. How can this be so? Imagine that a power plant for which we have very poor (highly uncertain) data is the primary source of CO₂ emissions in both life cycles. Imagine, too, that we have very good data on life cycle electricity consumption for both alternatives. In this case, we might simultaneously have the results that —

- a) life cycle emissions from alternative A are somewhere between 10 and 100 units, with 80% reliability;
- b) life cycle emissions from alternative B are somewhere between 20 and 130 units, with 80% reliability; and
- c) life cycle emissions from alternative A are virtually certain to be lower than those for alternative B.

In this example, we cannot use the individual results for each of the two alternatives and make an assessment of their difference assuming that the uncertainties are **independent** of each other. Instead, the uncertainties are highly correlated. This is because there is strong inter-dependence between the two results’ uncertainties — indeed, they stem primarily from the same uncertain data element within both life cycles.

More generally, this lack of independent uncertainty arises in LCA whenever there are one or more processes common to the two real-world life cycles being studied, about which we have uncertain data. Note that this is different from saying that we are using the same *process model* from our database in analyzing the two life cycles. The distinction may appear to be a subtle one, but it has critical implications for uncertainty analysis in LCA. This is the distinction between the model, and the modeled system, in LCA.

Reality, Viewpoint, and Model in LCA

Here we must point out that it is entirely reasonable — in fact, essential — to ask, “What exactly *are* the true ‘life cycle particulate emissions of a pound of newsprint’?” Not, “What do they equal, numerically?” but more deeply, “What do we *mean* by this concept?” Being right or wrong or uncertain about life cycle emissions implies that there *is* such a number “out there, in the real world.”

Central to most approaches for discussing and dealing with uncertainty is the notion of a “model.” The word “model” is an over-used one, and it gets applied in a confusing variety of ways. For our purposes, a **model** is a simplified characterization of some aspect of reality. Thus, there are really three levels to consider:

- 1) the “real world” — e.g., actual power plants, factories, and appliances in use;
- 2) some *aspects of*, or *viewpoints about*, the real world which we are interested in — e.g., the chain or tree of processes upstream or downstream of these power plants, factories, and appliances, as well as the inputs and outputs from each of these processes in the chain; and
- 3) a model of these aspects or this viewpoint — e.g., a life cycle inventory model.

It may help to briefly consider an analogy from outside of LCA. A dog is a complex, real-world creature. Now, there are many aspects to a dog — its behavior, its metabolism, its appearance, its feelings, etc. When the dog is beside us, we can look and see the true side-view appearance of the dog (at that moment in time). Of course, the dog is more complex than this view of him, however accurate it may be; he has many other features and aspects which are not captured in simply viewing him from the side. To complete the analogy, a painted picture of the side view of the dog is one step further removed from the full real-world complexity of the actual dog. The picture is a “model” of the side-view appearance of the dog. It attempts to capture many of the essential features of his side-view appearance.

Likewise with LCA. The real world contains mines and power plants and factories, dynamically exchanging inputs and outputs, and also consuming resources and emitting pollutants. A particular viewpoint or aspect of this reality is the (infinite) sequence of processes in the chain upstream of a given factory or power plant, and the sums of each of the flows released to the environment. Over a given time period, there are a certain number of inflows and outflows from each of the process sites in the chain. Given a selected method for co-product allocation, and a chosen time period for which the information is to pertain, there is a real-world LCI associated with a specific power plant’s electricity. Since the tree of processes to which this LCI pertains is infinite (and dynamic), we can never know this “true” LCI information. But it is this “true” set of induced pollution and resource flows upstream of the power plant that we are interested in when we build a model of the LCA viewpoint of this plant process.

Note that we are not considering the Life Cycle Assessment method itself to be the model of interest. Life Cycle Assessment is instead a particular way of characterizing the world, a series of conventions for viewing and describing some aspects of reality as a coherent and self-contained system. These conventions change as the field of LCA evolves — as users decide that there may be more useful and informative ways to characterize, compare, and evaluate the systems of interest.

Actual LCAs are then models, pictures painted (hopefully) in accordance with these viewing conventions. They are incomplete; they rely on imperfect and sometimes old information, not all of which pertains to the system we intend to be modeling — and these are the sources of uncertainty in our LCA models.

Error Types and Uncertainties in LCA

A new coal-fired electric power plant in Georgia is a real-world process, with inputs, a product (electricity), perhaps some co-products (e.g., fly ash), and various releases to the environment.

An LCA process model of this plant will be built from information obtained about this plant's inputs, outputs, and releases. This information will probably suffer from some **measurement/reporting error** — that is, the information obtained for the LCA may not be perfectly accurate. Also, the plant may continue to change over time, while our model of it will not (unless we continue to update it.) Thus, our model (data) will age relative to its subject, introducing **aging error**.

The same coal-fired electric power plant is also a member of many **classes**. More accurately, we should say that we can define many classes of which this plant is a member. For example —

- coal-fired electric power plants in Georgia
- new coal-fired plants in North America
- coal-fired plants in the southeast US
- electric power plants in the US

There are many other classes of which this plant is not a member, but with whose members it may nevertheless be expected to have some similarities with regard to the aspects we care about in LCA: inputs, coproducts, and releases per unit of product output. Our model of the Georgia plant might therefore be used as a surrogate for modeling such related classes, or individual members of them, in future LCAs. Such classes might include the following:

- coal-fired electric power plants in the southwest US
- new coal-fired plants in Canada

A feature of all the classes we have mentioned, and all the classes we are likely to be interested in within LCA, is that they have multiple members. In turn, these members are likely to exhibit **variability** in terms of the process level parameters: inputs, coproducts, and releases per unit of product output.

Instead of modeling a particular plant, we may wish to model a *class as a whole*. If we had data for all members of the class we could calculate the arithmetic mean process parameters for the class. Instead, we generally have data for only a subset of the class. If this subset were a randomly chosen *sample* of the class population, then we could use well-developed procedures to estimate the **sampling error** which arises from using the sample mean as an estimator or model of the class mean, for each parameter being modeled.

In LCA, however, we almost never have data for a randomly selected sample of the class. Generally the arithmetic mean parameters of the subset are adopted as estimates for the class, with some error. We suggest referring to this error as **subset error**, to distinguish it from the more readily quantified sampling error of random sampling.

Most generally, when we use data for (a subset of) one class to model another class (or a member of that other class), there is what Weidema and Wesnaes [1996] termed “technological distance” between the subject of the data and the object of our modeling. Actually, Weidema and Wesnaes used the term somewhat differently, referring to a “technological correlation between the data and the data quality goals.” As stated earlier, we feel that the correlation or relationship of importance for uncertainty analysis is that between the data subject and the modeled object.

Table 1 summarizes the principal factors, sources, and types of error contributing uncertainty to process-level estimates as a function of data usage. It also summarizes these same issues for tree-level estimates and life cycle level estimates, to which we now turn our discussion.

Table 1: Uncertainty of Process and Tree Level Data as a Function of Usage

Subject of Data	Modeled Object	Principal factors or error types contributing uncertainty to results
single process	same process	• measurement/reporting error (related to quality of source, carefulness, etc.) • aging error (temporal distance)
single process	different process in same class	• measurement/reporting error (related to quality of source, carefulness, etc.) • aging error (temporal distance) • technological distance between data subject and modeled object (influenced by in-class variability)
data for a subset within a class	average of same class	• measurement/reporting error (related to quality of source, carefulness, etc.) • aging error (temporal distance) • subset error (influenced by in-class variability, and size and representativeness of subset)
data for one or more processes	process in different class	• technological distance (subject vs. object) • temporal distance (subject vs. object)
data for one or more processes	average of different class	• technological distance (subject vs. object) • temporal distance (subject vs. object)
data for a tree	the tree	• errors in each process modeled • boundary truncation • unknown differences in method

Error Types and Uncertainties Related to Classes of Process Tree

A process tree model is an analytical construct using data for dozens, possibly hundreds, of separate processes. It starts with data for the driving processes, then it models how much of each input comes from which supplying processes, incorporates data for modeling those processes, models how much of each of their inputs comes from which supplying processes, and so on. It also incorporates methodological decisions about co-product allocation.

Now, starting from the actual process or class of processes in the real world which the driving process is meant to model, and given the same methodological decision about co-product allocation, *and* given a conceptual approach to characterizing each supplier linkage within the tree (which we will discuss further below), there is a real process tree about which our model is designed to help us draw conclusions. The issue of “tree data uncertainty” centres on this question: how uncertain are our LCI-based tree-level data and conclusions as predictors or estimators of the real-world life cycle inventories associated with the real-world tree, upstream of and including the driving process?

The most obvious way in which tree models and actual trees differ is in the model’s **boundary truncation**. The true tree, and therefore the true inventory of flows upstream of the driving process, includes the potentially non-negligible influence of dozens, maybe hundreds, of processes that are omitted from the model. The influence of this boundary truncation error cannot be known from the model results. It is obviously a function of the boundary truncation conventions used in the tree modeling. However, empirical studies are needed which test and report the results consequences of various boundary cut-off conventions upon total inventory results. An approach for doing such a study, and the results from preliminary investigations of such an approach, are reported in [Norris 2002].

We mentioned earlier that LCI models are simplified representations of a specific way of viewing the world. Both these models, and our real-world tree construct, require methodological decisions about co-product allocation and a conceptual approach to each supply linkage. Importantly, tree-level (e.g., cradle-to-gate) data in an LCI database embodies the decisions made on both of these issues. If the data are only cradle-to-gate summaries, then the user of these data has no ability to alter these decisions. Unfortunately, in some instances, the user cannot even know much about what these decisions were. In these instances, the user lacks important information about the *subject* of the original tree modeling. This limits the user's ability to take into account some potentially highly influential differences between the data's subject, and the object of the user's modeling. In turn, this increases the uncertainty the user has about the differences between his model results, and the real system inventory that the model is intended to be about.

There is another important consequence for results uncertainty of working with cradle-to-gate (that is, tree-level) data. This is an inability to take into account any correlations between results that are being used for a comparison. Recall from our earlier discussion that if one or more processes are modeled in both of two trees whose results are being compared, then the uncertainties in the two tree inventories are correlated. This correlation affects the reliability of (that is, the confidence we can have in) comparative conclusions about the two trees.

If there are positive correlations between two tree-level results uncertainties — caused by the same uncertain process being present in both systems — then assuming independence when assessing the reliability of comparative conclusions will tend to underestimate the reliability of such conclusions. If there were negative correlations between the two tree level results' uncertainties, then assuming independence would overestimate the confidence of comparative conclusions. This issue needs careful empirical study, to enable LCA practitioners and users to better understand the incidence and influence of positive (and possibly negative) correlations between tree-level results uncertainties.

Finally, let us consider briefly the issues surrounding the characterization of supply linkages, both in models and in the real-world conceptual systems being modeled.

Process models specify *how much* of each input is required by the process, per unit of product output. Trees further characterize *which* other processes provide each input to the using process. For a given real-world process, the supply of inputs can take many forms.

First, there may be one or multiple actual suppliers used for a given input. If multiple suppliers are used, and only one or a subset is modeled, this introduces uncertainty by increasing the potential technological distance between the model of the suppliers and the actual suppliers. But remember that models are representations of a viewpoint of reality. Before the uncertainty of a given model can be assessed, the LCA modeler needs to be clear about whether the system being modeled — the viewpoint — is of one or multiple suppliers. Model results *may* be about a true tree including only one of the suppliers.

Next, suppliers may be volatile (frequently changed) or stable. Typically, an LCA is intended to be representative for, or to capture the supply chain characteristics of, a given period in the order of a year. If suppliers are volatile, so that there are multiple suppliers within the year, then the multiple supplier issue described above is again relevant.

Finally, suppliers (and the supply linkage) may be **determinate** or **indeterminate**. An indeterminate supply linkage comes when the resource being purchased is pooled or well-mixed between supplying and using processes. The most important case is probably electricity that is purchased from the power pool ("grid"). Other important examples may be pipeline fuels, or commodities stockpiled by wholesalers or distributors. In such cases the actual supplying process is not knowable, and a mix should be considered as the only valid characterization of both the true system and its model.

Conclusions

In conclusion, we can see that uncertainty in LCA can neither be neglected, nor simply addressed. We see that information about uncertainty in life cycle inventory results cannot be fully captured within the LCI databases, because a significant share of this uncertainty arises in practice, based on the relationship between the data and the intended reality being modeled. Software and algorithms can and should be applied in the near term by researchers on a breadth of real-world case studies in order to more fully identify which are the major sources of uncertainty in life cycle assessment results. Once uncertainty types and sources have been prioritized, the LCA community can focus its efforts on characterizing them, reducing those which can be reduced, and structuring life cycle interpretation and results communication in a way that takes uncertainty into account.

References

Wilson, R.E, E. Crouch and L. Zeize, 1985. "Uncertainty in risk assessment." In Hoel D.G., R.A. Merrill and F.P. Perera, eds., *Risk quantification and regulatory policy*. Banbury Report # 19, Cold Spring Harbor Laboratories Press, Cold Spring Harbor, NY.

Funtowicz and Ravetz 1987: "Qualified Quantities – Towards and arithmetic of real experience." In *Measurement, Realism, and Objectivity*, J Forge (ed.), Dordrecht: Reidel, 59-88.

Funtowicz and Ravetz 1990, *Uncertainty and Quality in Science for Policy*. Dordrecht: Kluwer Academic.

Pate-Cornel, M.E., 1996. "Uncertainties in risk analysis: Six levels of treatment." *Reliability Engineering and System Safety*, 54: 95-111.

Cullen and Frey 1999: *Probabilistic Techniques in Exposure Assessment*. New York: Plenum

B P Weidema, M S Wesnaes (1996). Data quality management for life cycle inventories - an example of using data quality indicators. *Journal of Cleaner Production* 4(3-4):167-174.

EPA 1995. *Guidelines for assessing the quality of life-cycle inventory analysis*. US Environmental Protection Agency, EPA 530-R-95-010. April.

Research Triangle Institute 1992. *Assessing data quality for life cycle analysis*. Research Triangle Park, NC. August.

Norris, 2001: "Life Cycle Emission Distributions within the Economy: Implications for Life Cycle Impact Assessment", *Risk Analysis*, in press.