
A *k*-Nearest Neighbor Approach for Estimation of Single-Tree Biomass

Lutz Fehrmann¹ and Christoph Kleinn²

Abstract.—Allometric biomass models are typically site and species specific. They are mostly based on a low number of independent variables such as diameter at breast height and tree height. Because of relatively small datasets, their validity is limited to the set of conditions of the study, such as site conditions and diameter range. One challenge in the context of the current climate change discussion is to develop more general approaches for reliable biomass estimation. One alternative approach to widely used regression modelling are nonparametric techniques.

In this paper we use a *k*-Nearest Neighbor (*k*-NN) approach to estimate biomass for single trees and compare the results with commonly used regression models. The unknown target value of a certain tree is estimated according to its similarity to sample tree data stored in a database.

function of easily observable variables like such as diameter at breast height (d.b.h.) and tree height. Typically, these models are specific to the tree species and site conditions of the underlying particular study. Extrapolation beyond this set of particular conditions is critical.

Different attempts have been made to derive more general functions by meta-analyses of the published equations (e.g., Jenkins *et al.* 2003, Zianis and Mencuccini 2004, Chave 2005). In many cases such studies have been constrained by the absence of primary data and are focused on the reported regression functions only (Montagu *et al.* 2004). Therefore, one major goal of future research in the field of single-tree biomass estimation can be seen in the generalization of models based on compilation of empirical data from sample trees. Once a suitable single-tree database is given, nonparametric modelling approaches, such as the *k*-Nearest Neighbor (*k*-NN) method, might be suitable alternatives to regression modelling. The basic difference is that nonparametric models do not require concrete queries before they are developed.

Introduction

Estimation of forest biomass has gained importance in the context of the legally accepted framework of the United Nations Framework Convention on Climate Change and the Kyoto Protocol. Reliable and general estimation approaches for carbon sequestration in forest ecosystems are needed (Brown 2001, Joosten *et al.* 2003, Rosenbaum *et al.* 2004, Wirth *et al.* 2003). In the past, the standard methodology in single-tree biomass estimation was based on fitting parametric regression models with relatively small datasets. Numerous models have been built from destructive sampling studies, most of which are allometric functions. They allow predicting tree biomass as a

Methods

k-NN Technique

The *k*-NN approach is a nonparametric and instance-based machine learning algorithm. It is known as one of the oldest and simplest learning techniques based on pattern recognition and classification of unknown objects. It was described as a nonparametric approach for discriminant analysis (lazy similarity learning algorithm) by Fix and Hodges (1989) or Cover and Hart (1967), for example.

This approach classifies an unknown feature of an object (an instance) based on its “overall” similarity to other known

¹ Institute of Forest Management, Georg-August-Universität Göttingen, Büsgenweg 5, D-37077 Göttingen. E-mail: lfehrma@uni-forst.gwdg.de.

² Professor of Forest Assessment and Remote Sensing, Institute of Forest Management, Georg-August-Universität Göttingen, Büsgenweg 5, D-37077 Göttingen. E-mail: ckleinn@gwdg.de.

objects. Therefore, the instances with known target values are stored in a database (the so-called training data).

To estimate the unknown feature of a query instance, the most similar known instances are identified by means of a set of known variables. The weighted or unweighted mean of the target variable of a number k of nearest instances (neighbors) to the unknown instance is then assigned. To identify the most similar training instances, it is necessary to define measures of similarity and quantify their distance or dissimilarity to the query instance (Haendel 2003).

In contrast to parametric models, the result of the k -NN estimation is not a “global function” for the entire feature space, but a local approximation of the target value that changes in every point of the feature space depending on the nearest neighbours that can be found for a certain query point (Mitchell 1997).

In forestry, applications of this approach can be found in Haara *et al.* (1997), Korhonen and Kangas (1997), Maltamo and Kangas (1998), Niggemeyer (1999), Tommola *et al.* (1999), and Hessenmöller (2001). In this paper, the methodology is mainly used to estimate stand parameters or as an alternative to parametric growth models. Sironen *et al.* (2003) applied a k -NN approach for growth estimations on single-tree data. Applications of different nonparametric approaches including k -NN are also in Malinen (2003a, 2003b), Malinen and Maltamo (2003), and Malinen *et al.* (2003).

The k -NN technique has long proved applicable and useful in the context of integration of satellite imagery into large-scale forest inventories estimations (Moer and Stange 1995, Tomppo 1991). Satellite images are classified using the similarity of spectral signatures of single-pixel values (Holmström *et al.* 2001, McRoberts *et al.* 2002, Stürmer and Köhl 2005).

For local approximation of a continuous target value, the k -NN algorithm assigns the mean of the target values of a certain number of most similar training instances to the query instance as

$$\hat{f}(x_q) \leftarrow \frac{\sum_{i=1}^k f(x_i)}{k} \quad (1)$$

Where:

$\hat{f}(x_q)$ is the estimator for the unknown target value of a query instance x_q .

$f(x_i)$ are the known target values of training instances.

k is the number of nearest neighbours used for estimation.

To quantify the dissimilarity between instances and to identify a number of k nearest neighbours, known measures of proximity from multivariate analyses, such as discriminant or cluster analysis, may be used. For practical application the Minkowski metric or L-norm is a suitable and flexible multivariate distance measure (Bortz 1989, Backhaus *et al.* 1996):

$$d_{i,j} = \left[\sum_{r=1}^n |x_{ir} - x_{jr}|^c \right]^{\frac{1}{c}} \quad (2)$$

where:

$d_{i,j}$ = the distance between two instances i and j , x_{ir} and x_{jr} being the values of the r^{th} variable for the respective instance.

n = the number of considered variables.

$c \geq 1$ = the Minkowski constant.

In case of $c = 1$, the result of this metric is the so called Manhattan or taxi driver distance, which is the sum of all variables differences. For $c = 2$, this measure is the Euclidean distance in an n -dimensional feature space.

To take the unequal importance of different variables for the development of the target value into account and to avoid the distorting influence of different scaled feature spaces, the variables have to be standardized and weighted according to their influence. Because the single variable distances are explicitly obvious in the given distance metric, the standardization and weighting can be included in the calculation of an overall distance by modifying it to the following:

$$dw_{i,j} = \left[\sum_{r=1}^n \left(w_r \frac{|x_{ir} - x_{jr}|}{\delta_r} \right)^c \right]^{\frac{1}{c}} \quad (3)$$

where:

$dw_{i,j}$ = the weighted distance.

w_r = the weighting factor for variable r .

δ_r = a standardization factor that can be coupled to the range of the variable.

In our study we set δ_r to $2\sigma_r$ whereas σ is the standard deviation of the respective variable.

Even if both steps are a kind of transformation of the feature spaces of the considered variables, one should distinguish between feature standardization and weighting. While standardization is necessary to ensure the comparability of the single variable distances, the weighting of the different variables in a multidimensional space is an expression of their unequal relevance for the target value (Aha 1998, Wettschereck 1995).

Feature weighting can have a great influence on the identification of the nearest neighbours, making it relevant for the quality of the derived estimation. Suitable weighting factors can be derived from several alternatives. Tomppo *et al.* (1999) proposes deriving feature weights based on the coefficient of correlation between the different variables and the target value. Another possibility is using the relation between the regression coefficients of the included variables from a suitable regression model to derive the weighting factors. Iterative optimization algorithms such as genetic algorithm or simulated annealing can also be used to find an appropriate relation of feature weighting factors (Tomppo and Halme 2004).

If the distances between a query point and all training instances in the database are known and the k nearest neighbours are identified, they can also be used to derive a weighted mean. According to Sironen *et al.* (2003), or similarly Maltamo and Kangas (1998), the weighting of the neighbours according to their distance can thereby be derived as

$$w_k = \frac{\left(\frac{1}{d_{q,i}}\right)^t}{\sum_{i=1}^k \left(\frac{1}{d_{q,i}}\right)^t} \quad (4)$$

where:

w_k = the weight of the k^{th} neighbour.

$d_{q,i}$ = the distance between a query point x_q and the neighbour x_i .

t = a weighting parameter that influences the kernel function.

Implementing this distance-weighted mean as estimator formula (1) becomes

$$\hat{f}_w(x_q) \leftarrow \frac{\sum_{i=1}^k w_k f(x_i)}{\sum_{i=1}^k w_k} \quad (5)$$

In this case, the estimator is equivalent to the *Nadaraya-Watson estimator* (Atkeson *et al.* 1997, Haendel 2003, Nadaraya 1964, Watson 1964). Because of the decreasing influence of training instances with increasing distance, all training instances can be included in the estimation process in this approach, which is also known as Shepard's method (Shepard 1968).

Even if the k -NN algorithm is referred to as a nonparametric method in the context of searching a number of nearest neighbours, this description does not apply for the distance function that is used. In the basic k -NN approach, the weighting factors for the different variables, which are normally defined in a deterministic manner, and the parameters k , n , c , δ and t of the above mentioned distance function (3) and estimator (5) are defined globally.

As a result of an asymmetric neighbourhood at the extremes of the distribution of observations, instance-based methods come with a typical bias-variance dilemma. The number of neighbours considered in the estimation must be determined as a compromise between an increasing bias and the decreasing variance of estimates with an increasing number of neighbours (Katila 2004). To find an approximation for an optimal number for k , we applied the root mean square error (rMSE%) as error criterion. The objective criteria is the minimization of the rMSE% by means of a leave-one-out cross validation with a changing size of the considered neighbourhood and/or the parameter setting in the distance and weighting function. In a cross validation, a query instance is a tree that is excluded from the training instances and for which estimation is derived based on the $N-1$ remaining trees. Each training instance is in turn used as query instance (Malinen *et al.* 2003). The rMSE% is then calculated as

$$rMSE\% = 100 \times \frac{\sqrt{\frac{1}{n} \sum_{i=1}^n (x_{ir} - \hat{x}_{ir})^2}}{\bar{\hat{x}}_r} \quad (6)$$

where:

x_{ir} = the observed value of variable r for instance i .

$\hat{x}_{ir}$ = the respective estimated value.

n = the number of observations.

$\bar{\hat{x}}_r$ = the mean of estimates for the target variable r .

Because of the high number of possible parameter combinations, the iterative process was reduced to about 50 different combinations in which the determination of the starting values for the feature weights were based on expertise obtained from the relation of regression coefficients as result of a regression analysis with the respective variables.

Reference Model

As reference to the k -NN estimations we derived an allometric regression model based on the same dataset. Independent variables are dbh and tree height. To consider for inherent heteroscedasticity, an ordinary least square (OLS) regression was built with log transformed variables and aboveground biomass (agb) as the dependent variable (Sprugel 1983). The estimated regression coefficients are shown in table 1.

Data

To evaluate the k -NN approach in comparison to parametric regression models, we built a single-tree biomass database with training instances from various destructive biomass studies. In this study we used a subset of $N = 323$ Norway spruce trees (*Picea abies* [L.] Karst.) that were compiled from different publications and datasets from central Europe. Parts of the database come from a study of Wirth *et al.* (2003). Additional datasets were taken from literature and project reports.

Results

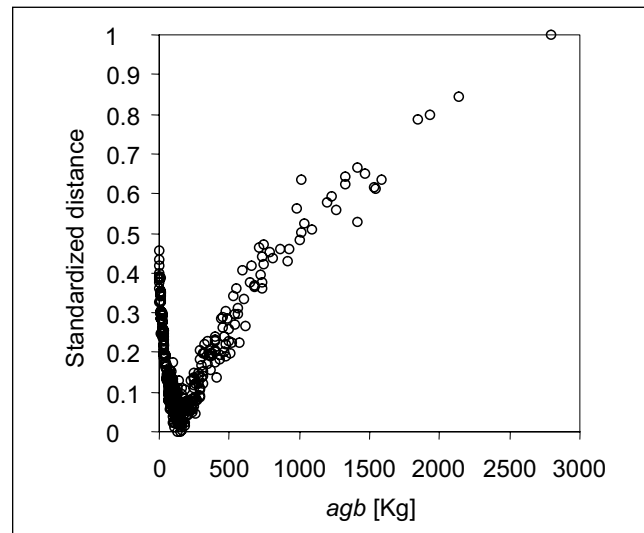
Calculating the distance between a query point and all training data leads to a certain order of instances according to

their similarity to the unknown query instance (fig. 1). For the given example, at first only the variables dbh and tree height were used in the distance function. Using a nearest neighbour bandwidth, the distance to which neighbours are considered for the estimation is set to the distance of the k^{th} neighbour (Atkeson *et al.* 1997).

The disadvantage of this approach is that a fixed bandwidth selection may increase bias as result of the asymmetric neighbourhood in the extremes of the feature space. Nonparametric approaches such as the k -NN method are known to be inappropriate for any kind of extrapolations. In addition, a certain edge effect exists within the feature range of the training data. This fact makes it difficult to compare the k -NN based estimations with a given regression model. Figure 2 shows the effect of increasing the neighbourhood, especially on estimations for the biggest trees in this dataset.

The smaller size of the considered neighbourhood leads to a lower bias at the extremes of the feature space. At the same

Figure 1.—Training data ordered according to their distance and their respective target values (agb) for a given query point.



agb = aboveground biomass.

Table 1.—Estimated coefficients, R^2 , and residual standard error for the allometric reference model. Dependent variable is agb in kilogram dry mass.

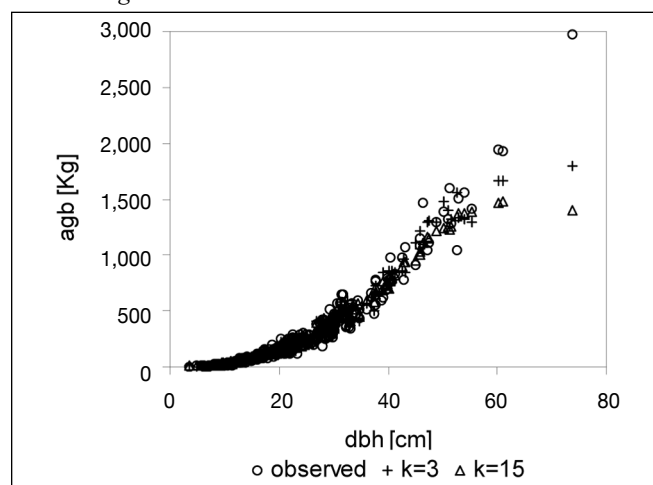
Model formulation	Linearized form	$\ln(a)$	b	c	R^2	Residual standard error
$agb = a \cdot dbh^b \cdot h^c$	$\ln(agb) = \ln(a) + b \cdot \ln(dbh) + c \cdot \ln(h)$	-2.651	1.888	0.699	0.98	0.1864

agb = aboveground biomass.

time, however variance is obviously increasing because fewer values of the target values are averaged.

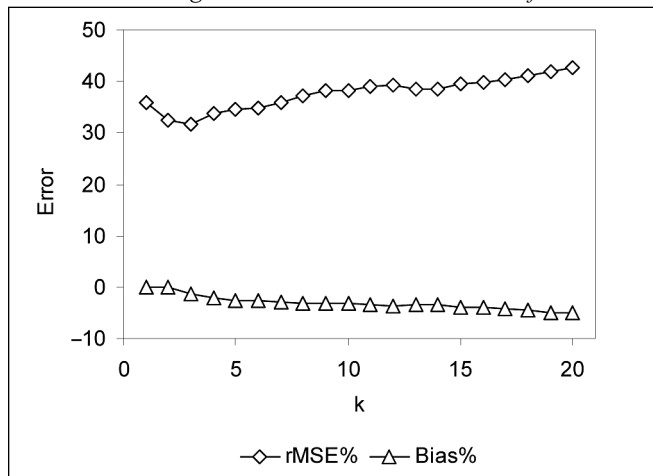
The resulting rMSE% values for different sizes of the neighborhood were calculated by means of a leave-one-out cross validation of the whole dataset for different values of k . Figure 3 shows that in case of the underlying data and the used variables, a minimum error can be found for three neighbors. It must be considered that this optimum size of the neighborhood is only valid for the given parameter setting and this certain

Figure 2.—Observed agb and estimations based on $k = 3$ and $k = 15$ neighbors.



agb = aboveground biomass.

Figure 3.—rMSE % and Bias% of the k -NN estimation calculated in a leave-one-out cross validation for different sizes of the considered neighbourhood. In this case only the variables d.b.h. and tree height were included in the distance function.



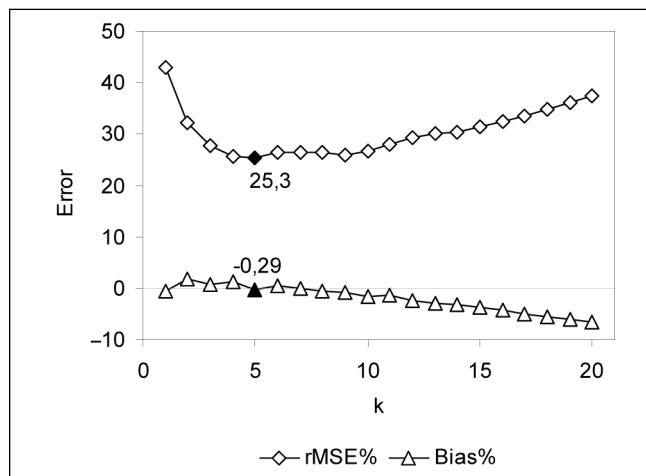
d.b.h. = diameter at breast height.

dataset. The respective rMSE% calculated for the adapted reference model was about 19 percent lower than that of the k -NN estimation.

A lower error can be achieved by integrating further single-tree variables, such as tree age or crown length. Figure 4 shows an example where these additional variables were included in the distance calculation. It is obvious that the optimal number of neighbours changes in this case to five. The amount of available training data with information for these search variables decreased to 181 trees.

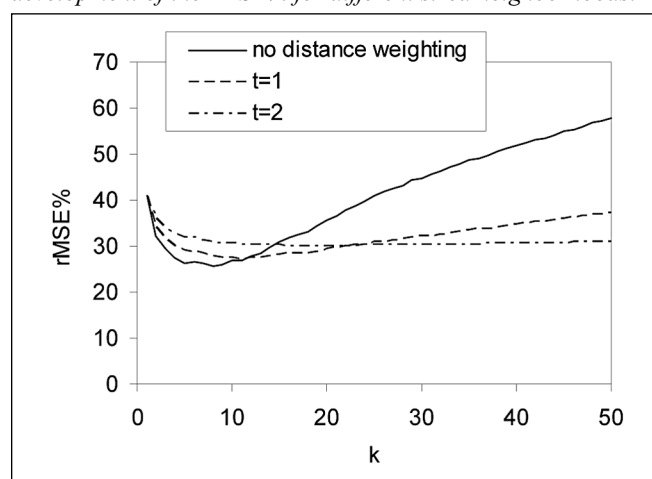
One possibility to lower the influence of the systematic error is to use a kernel function that attenuates the influence of neighbours according to their increasing distance. The distance-weighting function we used in this approach can be modified by changing the parameter t . As figure 5 shows, we achieved the lowest errors without any distance weighting in this case. Different simulations with other tree species and different combinations of variables showed that the influence of the parameter t on the error is highly dependant on the underlying dataset and the choice of search variables. Similar to the determination of the parameter c that influences the type of distance metric as well as the relation of feature-weighting factors, the optimum parameter setting differs highly between different datasets.

Figure 4.—rMSE% and Bias% for a certain parameter setting using the variables d.b.h., tree height, tree age, and crown length for the distance function.



d.b.h. = diameter at breast height.

Figure 5.—Influence of the distance-weighting parameter t on the development of the $rMSE\%$ for different sized neighborhoods.



Discussion

In the given example, the errors of the k -NN estimations were higher than those of an allometric regression model derived from the same dataset and with the same independent variables. One reason for the errors is that we only used one combination of feature weights and/or parameter settings. The goal of this study was primarily to evaluate the general applicability of the k -NN method for single-tree biomass estimation. Future work will be focused on optimization of this method for the given purpose.

Another reason for the comparatively bad performance of the approach in comparison to the given regression model is that the number of training data as well as the number of variables included in the distance function was untypically low for a nonparametric method. One advantage of the k -NN method is that it is easily possible to include a high number of independent search variables in the distance function. This advantage was not used in this basic and general example. More variables and information components can be included in the estimation process. For examples, the process could include site parameters such as the height above sea level, the site quality, geographic coordinates, or further information on tree species. If enough variables with a certain discriminatory

power in the context of dissociating trees of different species are available and are able to bring the training data in a correct order according to their distances, estimations over different species are possible.

For the optimization of the k -NN estimations, the high number of parameters for the distance and weighting function we used in this example causes problems. An approximation for an optimal combination or relation of parameters can be found by using iterative processes such as optimization algorithms. The target in this case is the minimization of the error criterion (e.g., the $rMSE\%$) by a stepwise change of the parameter settings. For the given example we only used a low number of iterations, whereas the starting values were predefined based on expertise gained in the regression analysis and the first experience with a software application of the k -NN algorithm. It must be assumed that the given intermediate results are far from an optimal solution and that future work can enhance the performance of the k -NN method which might then be an alternative to the given approaches, particularly in the context of the generalization of biomass models.

Conclusions

Trees can be interpreted as instances of more or less one basic form, consisting of an individual pattern of principal components such as stem, branches, leaves, and roots. If certain key variables on a single-tree level are known, pattern recognition algorithms such as the k -NN method can be used to identify the most similar instances (trees) from a database and use their known target values to derive estimations for unknown instances. Often additional meta-information about forest stands, site characteristics, or species-specific information such as mean wood density are available and can be used in the context of this methodology. Different authors (Hessenmöller 2001, Malinen 2003a, Sironen *et al.* 2003) have proved this nonparametric method applicable and useful as well for single-tree applications, in which its performance is obviously highly dependent on the amount of training data. One of the main future challenges in the field of biomass estimation will be the generalization of estimation approaches and/or models.

To ensure a certain reliability of the generalized models, the compilation of destructive sampled data will be necessary. If it is possible to implement a single-tree database in a central and free accessible place, instance-based methods can also be applied as server-client applications in future.

Acknowledgments

The German Science Foundation supports this research with research grant KL594/4. The authors highly appreciated this support. A project such as this crucially depends on colleague researchers making some of their data available; this is a notoriously difficult undertaking. Therefore, the authors are sincerely indebted to C. Wirth, Dr. A. Akça, and Dr. A. Mench for providing their biomass raw data that considerably helped to build a reasonably sized database. The authors also highly appreciated Lars Hinrichs' critical review of the manuscript.

Literature Cited

- Aha, D.W. 1998. Feature weighting for lazy learning algorithms. Unpublished report. Washington, DC: Navy Center for Applied Research in Artificial Intelligence. 20 p.
- Atkeson, C.G.; Moore, A.W.; Schaal, S. 1997. Locally weighted learning. *Artificial Intelligence Review*. 11(1-5): 11-73.
- Backhaus, K.; Erichson, B.; Plinke, W.; Weiber, R. 1996. Multivariate Analysemethoden. Eine anwendungsorientierte Einführung. 8. Auflage: Springer-Verlag. 591 p. In German.
- Bortz, J. 1989. Statistik für Sozialwissenschaftler. 3. Auflage: Springer-Verlag. 900 p. In German.
- Brown, S. 2001. Measuring carbon in forests: current status and future challenges. *Environmental Pollution*. 116(3): 363-372.
- Chave, J.; Andalo, C.; Brown, S.; Cairns, M.A.; Chambers, J.Q.; Eamus, D.; Fölster, H.; Fromard, F.; Higuchi, N.; Kira, T.; Lescure, J.-P.; Nelson, B.W.; Ogawa, H.; Puig, H.; Riéra, B.; Yamakura, T. 2005. Tree allometry and improved estimation of carbon stocks and balance in tropical forests. *Oecologia*. 145(1): 87-99.
- Cover, T.; Hart, P. 1967. Nearest neighbour pattern classification. *IEEE Transactions on Information Theory*. 13(1): 21-27.
- Fix, E.; Hodges, J.L., Jr. 1989. Discriminatory analysis—nonparametric discrimination: consistency properties. *International Statistical Review*. 57: 238-247.
- Haara, A.; Maltamo, M.; Tokola, T. 1997. The k -nearest neighbour method for estimating basal-area distribution. *Scandinavian Journal of Forest Research*. 12: 200-208.
- Haendel, L. 2003. Clusterverfahren zur datenbasierten Generierung interpretierbarer Regeln unter Verwendung lokaler Entscheidungskriterien. Dissertation an der Fakultät für Elektrotechnik und Informationstechnik der Universität Dortmund. 120 p. In German.
- Hessenmöller, D. 2001. Modelle zur Wachstums- und Durchforstungssimulation im Göttinger Kalkbuchenwald. Dissertation zur Erlangung des Doktorgrades der Fakultät für Forstwissenschaften und Waldökologie der Georg-August Universität Göttingen. Logos Verlag Berlin. 163 p. In German.
- Holmström, H.; Nilsson, M.; Ståhl, G. 2001. Simultaneous estimations of forest parameters using aerial photograph interpreted data and the k -nearest neighbour method. *Scandinavian Journal of Forest Research*. 16(1): 67-78.
- Jenkins, J.C.; Chojnacky, D.C.; Heath, L.S.; Birdsey, R.A. 2003. National-scale biomass estimators for United States tree species. *Forest Science*. 49(1): 12-35.

-
- Joosten, R.; Schumacher, J.; Wirth, C.; Schulte, A. 2003. Evaluating tree carbon predictions for beech (*Fagus sylvatica* L.) in western Germany. *Forest Ecology Management*. 189: 87-96.
- Katila, M. 2004. Error variations at the pixel level in the *k*-nearest neighbour predictions of the Finnish multi-source forest inventory. *Proceedings, first GIS & remote sensing days*. 408 p.
- Korhonen, K.T.; Kangas, A. 1997. Application of nearest-neighbour regression for generalizing sample tree information. *Scandinavian Journal of Forest Research*. 12: 97-101.
- Malinen, J. 2003a. Prediction of characteristics of marked stand and metrics for similarity of log distribution for wood procurement management. Finland: University of Johensuu. Dissertation.
- Malinen, J. 2003b. Locally adaptable non-parametric methods for estimating stand characteristics for wood procurement planning. *Siva Fennica*. 37(1): 109-118.
- Malinen, J.; Maltamo, M.; Harstela, P. 2003. **Application of most similar neighbour inference for estimating marked stand characteristics using harvester and inventory generated stem database.** *International Journal of Forest Engineering*. 33.
- Malinen, J.; Maltamo, M.; Verkasalo, E. 2003. **Predicting the internal quality and value of Norway spruce trees by using two non-parametric nearest neighbour methods.** *Forest Products Journal*. 53(4): 85-94.
- Maltamo, M.; Kangas, A. 1998. Methods based on *k*-nearest neighbour regression in the prediction of basal area diameter distribution. *Canadian Journal of Forest Research*. 28: 1107-1115.
- McRoberts, R.; Nelson, M.D.; Wendt, D.G. 2002. **Stratified estimation of forest area using satellite imagery, inventory data, and the *k*-nearest neighbour technique.** *Remote Sensing of Environment*. 82: 457-468.
- Mitchell, T. 1997. *Machine learning*. New York: McGraw-Hill. 432 p.
- Moer, M.; Stange, A.R. 1995. **Most similar neighbour: an improved sampling inference procedure for natural resource planning.** *Forest Science*. 41: 337-359.
- Montagu, K.D.; Düttmer, K.; Barton, C.V.M.; Cowie, A.L. 2004. **Developing general allometric relationships for regional estimates of carbon sequestration—an example using *Eucalyptus pilularis* from seven contrasting sites.** *Forest Ecology Management*. 204: 113-127.
- Nadaraya, E.A. 1964. On estimating regression. *Theory of Probability and its Applications*. 9: 141-142.
- Niggemeyer, P.; Schmidt, M. 1999. Estimation of the diameter distributions using the *k*-nearest neighbour method. In: Pukkala, T.; Eerikäinen, K., eds. *Growth and yield modelling of tree plantations in South and East Africa*. University of Johensuu, Faculty of Forestry Research Notes. 97: 195-209.
- Rosenbaum, K.L.; Schoene, D.; Mekouar, A. 2004. **Climate change and the forest sector: possible national and subnational legislation.** *FAO Forestry Paper 144*. New York: United Nations, Food and Agriculture Organization. 73 p.
- Shepard, D. 1968. A two-dimensional interpolation function for irregularly spaced data. *Proceedings, 23rd National Conference of the ACM*.
- Sironen, S.; Kangas, A.; Maltamo, M.; Kangas, J. 2003. **Estimating individual tree growth with nonparametric methods.** *Canadian Journal of Forestry Research*. 33: 444-449.
- Sprugel, D.G. 1983. Correcting for bias in log-transformed allometric equations. *Ecology*. 64: 209-210.

-
- Stümer, W.; Köhl, M. 2005. Kombination von terrestrischen Aufnahmen und Fernerkundungsdaten mit Hilfe der *k*-Nächste-Nachbarn-Methode zur Klassifizierung und Kartierung von Wäldern. *Fotogrammetrie Fernerkundung Geoinformation*. 1/200: 23-36. In German.
- Tommola, M.; Tynkkyen, M.; Lemmetty, J.; Harstela, P.; Sikanen, L. 1999. Estimating the characteristics of a marked stand using *k*-nearest-neighbour regression. *Journal of Forest Engineering*. 10(2): 75-81.
- Tomppo, E. 1991. Satellite imagery-based national inventory of Finland. *International Archives of Photogrammetry and Remote Sensing*. 28(7-1): 419-424.
- Tomppo, E.; Goulding, C.; Katila, M. 1999. Adapting Finnish multi-source forest inventory techniques to the New Zealand preharvest inventory. *Scandinavian Journal of Forest Research*. 1: 182-192.
- Tomppo, E.; Halme, M. 2004. Using coarse scale forest variables as ancillary information and weighting of variables in *k*-NN estimation: a genetic algorithm approach. *Remote Sensing of Environment*. 92: 1-20.
- Wettschereck, D.; Aha, D.W. 1995. **Weighting features**. In: Voloso, M.; Aamodt, A., eds. 1995. *Proceedings, first international conference on case-based reasoning*. Sesimbra, Portugal: Springer: 347-358.
- Wirth, C.; Schumacher, J.; Schulze, E.D. 2003. **Generic biomass functions for Norway spruce in Central Europe—a meta analysis approach toward prediction and uncertainty estimation**. *Tree Physiology*. 24: 121-139.
- Zianis, D.; Mencuccini, M. 2004. On simplifying allometric analyses of forest biomass. *Forest Ecology Management*. 187: 311-332.