

Designing plots for precise estimation of forest attributes in landscapes and forests of varying heterogeneity

Andrew J. Lister and Laura P. Leites

Abstract: Models of relationships among forest inventory sampling efficiency and cluster plot configuration variables inform decisions by inventory planners. However, relationships vary under different spatial heterogeneity scenarios. To improve understanding of how spatial patterns of forests affect these relationships, we implemented a factorial experiment by simulating forest pattern at both the landscape and stand scales. We sampled these simulated forests with a variety of cluster plot configurations, calculated coefficient of variation (CV) of trees per hectare for each replicate, and tested the relationships among CV and the heterogeneity and cluster plot configuration factors within a linear mixed model framework. Both landscape- and stand-scale pattern aggregation had a significant relationship with CV. Changing cluster plot configuration factors did little to change the overall CV when using larger subplots but had some important effects when using smaller subplots. These impacts were stronger in the more uniform landscapes. Results were opposite for stand-scale heterogeneity; changing plot configuration in areas with aggregated patterns had a stronger impact than it did in areas with more uniform patterns. Results of this study reveal the importance of accounting for spatial pattern at multiple scales when making cluster configuration choices if the goal is statistical efficiency.

Key words: cluster plot design, forest inventory design optimization, forest inventory efficiency, forest pattern simulation, forest sampling simulation.

Résumé : Les modèles qui relient l'efficacité d'échantillonnage de l'inventaire forestier et les variables de configuration des grappes de parcelles éclairent les décisions des planificateurs d'inventaire. Cependant, ces relations varient en fonction des différents scénarios d'hétérogénéité spatiale. Afin d'améliorer la compréhension de la façon dont les patrons spatiaux des forêts influencent ces relations, nous avons mis en œuvre une expérience factorielle en simulant les patrons forestiers aux échelles du paysage et du peuplement. Nous avons échantillonné ces forêts simulées avec une variété de configurations de grappes de parcelles, calculé le coefficient de variation (CV) du nombre d'arbres à l'hectare pour chaque répétition, et testé les relations entre le CV et les facteurs d'hétérogénéité spatiale et de configuration des grappes de parcelles, dans un cadre de modèle linéaire mixte. Les patrons d'aggrégation spatiale avaient une relation significative avec le CV, tant à l'échelle du paysage que du peuplement. La modification des facteurs de configuration des grappes de parcelles n'a pas changé le CV global lors de l'utilisation de sous-parcelles plus grandes, mais a eu des effets importants lors de l'utilisation de sous-parcelles plus petites. Ces effets étaient plus importants dans les paysages plus uniformes. Les résultats étaient à l'opposé pour l'hétérogénéité à l'échelle du peuplement; la modification de la configuration des parcelles dans les zones avec des patrons agrégés a eu un impact plus marqué que dans les zones avec des patrons plus uniformes. Les résultats de cette étude révèlent l'importance de tenir compte du patron spatial à plusieurs échelles lors des choix de configuration des grappes si l'objectif est l'efficacité statistique. [Traduit par la Rédaction]

Mots-clés : configuration des grappes de parcelles, optimisation de la conception de l'inventaire forestier, efficacité de l'inventaire forestier, simulation de patrons forestiers, simulation d'échantillonnage forestier.

Introduction

The impacts of the spatial distribution of forest resources on the efficiency of forest monitoring systems vary with the scale of analysis and with the attribute of interest. Spatial variation exists at both the landscape (Heilman et al. 2002; Haddad et al. 2015) and local (Stoyan and Penttinen 2000) scales and is due to a combination of human and natural biotic and abiotic influences. This wide range of variability at multiple scales presents a challenge to planners seeking efficient forest monitoring system designs, particularly as developing countries create forest monitoring systems that meet requirements for participation in degradation and

deforestation reduction incentive programs like the United Nations' Reducing Emissions from Deforestation and Forest Degradation (REDD) program (UNFCCC 2016).

Monitoring can generally be made more efficient by sampling as opposed to exhaustive measurement of the resource of interest. Sampling design is the most consequential decision with regard to efficiency, as it involves decisions related to field plot design, inferential paradigm, form of the estimator, number of plots, sample unit selection process, and data collection protocols (Thompson 2012, p. 2). An important design decision is whether to incorporate auxiliary information in the form of remote sensing data. For example, remote sensing imagery can help in planning and other

Received 1 December 2020. Accepted 31 March 2021.

A.J. Lister. USDA Forest Service, Northern Research Station, Forest Inventory and Analysis, 3460 Industrial Drive, York, PA 17402, USA.

L.P. Leites. Department of Ecosystem Science and Management, Penn State University, 312 Forest Resources Building, University Park, PA 16802, USA.

Corresponding author: Andrew Lister (email: andrew.lister@usda.gov).

© 2021 The Author(s). Permission for reuse (free in most cases) can be obtained from copyright.com.

logistical aspects of the inventory, or it can be used in estimation and inference, such as in the case of model-assisted or model-based inference. Traditional design-based inference is based on probabilistic sample unit selection and the distribution of all possible estimates obtainable using a given sampling protocol; remote sensing data can be used to “assist” inference in this case by providing input to a model that describes the population while still relying upon the probabilistic nature of the sampling design to make inferences. Using remote sensing data in a model-based paradigm, on the other hand, entails reliance on a model for inference about population parameters (Gregoire 1998; McRoberts 2010). In the context of this study, we are using a finite population sampling paradigm and employ design-based inference without the use of auxiliary data. This is commonly done in many large area forest inventories around the world, assuming that the use of stratification is not considered a model-assisted technique.

One approach to forest inventory design is to seek the best precision for a fixed number of plots, i.e., to reduce the variance of the estimate of a mean or total, which is often referred to as sampling error or relative standard error (Thompson 2012). Coefficient of variation (CV), which is a component of sampling error, is a commonly used index of the variability of a sample. Plot configuration (hereafter, plot design) has direct influence on the CV of an attribute by affecting the average deviation between each plot value and the mean plot value; if a plot is configured such that each plot is a microcosm of the population as a whole, the average deviation will be small, CV will be small, and precision will be improved for a fixed number of plots. It is also common to design a forest inventory such that estimates will meet an allowable error (AE) criterion, such as a sampling error of 10% of the estimate at the 95% confidence level (IPCC 2006). Plot design thus has an indirect effect on required sample size needed to achieve AE through its effects on CV (eq. 1):

$$(1) \quad n_{\text{req}} = \left(\frac{CV\% \times t}{AE\%} \right)^2$$

where n_{req} is the required sample size to achieve a specified AE% given a known or hypothesized CV% of the attribute of interest, and t is the Student's t value associated with the desired confidence level (Loetsch and Haller 1973).

Plot design factors that can be altered to adjust CV include size and shape for single subplot designs, as well as count and separation distance for multiple subplot or cluster designs. In cluster designs, primary units (sampling units, often referred to simply as plots) are composed of more than one secondary unit (measurement units, often simply referred to as subplots) distributed in a defined pattern such as a line, cross, or L-shape. For a fixed number of plots, separation of subplots in space leads to plot-level estimates that more closely resemble the sample mean by avoiding redundant sampling effort in neighbouring patches of land with similar conditions.

In general, plots acquire a large amount of new information over short distances as plot area increases, leading to sharp improvements in precision. As plot area increases beyond a certain threshold, however, there are diminishing improvements with added area because each plot's status comes to resemble that of the sample mean; this phenomenon is depicted conceptually in Supplementary Material S1¹. Smith (1938) was among the first to identify this negative exponential relationship between plot area and variance, and since then, it has been explored in a forestry context by several authors, as reviewed by Lynch (2017). For a given plot design, the magnitude of the absolute value of the exponent depicted in Supplementary Material S1¹ varies by spatial pattern

scenario; Lynch (2017) reports values in the literature ranging between approximately -0.1 (for aggregated patterns or those with a spatial trend, e.g., Reich and Arvanitis (1992)) to approximately -1 (for a completely random pattern, per the discussion in Zeide (1980)). Each plot design factor affects precision in a similar manner, but the interactions among these, landscape- and local-level heterogeneity, and inventory precision are difficult to predict without a systematically planned analysis approach like a designed experiment.

There is thus a lack of clear guidelines for how the interaction of multiple levels of heterogeneity affect cluster plot design decisions; this can lead to decisions with costly repercussions. For example, the United Nations Food and Agriculture Organization at one time suggested cluster plot designs with four 0.5 ha subplots separated by 500 m in a square configuration (Branthomme 2004). Some national forest inventories employ much smaller clusters, like that of the United States, which uses four 0.17 ha subplots, separated by 37 m and arranged in a triangular pattern (Bechtold et al. 2005). Other inventories, like those for nonforest trees that occur in sparse clumps, are conducted with larger subplots or with linear transects and line intercept sampling (Kleinn et al. 2001). With such a broad range of plot design choices, it is critical that there exist conceptual and empirical models that provide heuristics to help guide inventory planners with plot design and configuration decisions. With an improved understanding of which design variables are important, and how their importance responds to different spatial heterogeneity scenarios, inventory designers can improve decisions.

Many studies select plot designs using information from stem-mapped stands that are purposively chosen and deliberately restricted to forested areas (e.g., Schreuder et al. 1987; Picard et al. 2018). Others have chosen to extract subsets of trees using different plot designs from existing forested inventory plots (Lynch 2003; Picard et al. 2004; Yim et al. 2015). Still others have used models to create artificial forests and then simulated point or other types of sampling (Arvanitis and O'Regan 1967; Mackisack and Wood 1990; Brink and Schreuder 1992; Hou et al. 2015; Gove 2017). None, to our knowledge, have used a designed experiment to model how effects of different scales of heterogeneity interact with each other and with several plot design factors to reduce variance of forest inventory estimates. Understanding these interactions is critical when achieving an AE is the goal, or when paired with cost estimates. However, there is often a lack of meaningful cost data to guide decisions due to uncertainties about field or other logistics. In such cases, having knowledge of the relationships among design factors, the population's spatial structure, and n_{req} is valuable. It helps planners identify which plot design components are most impactful and suggests ranges of values for plot design variables.

We conducted a factorial simulation experiment to investigate the precision impacts of these interactions. We tested the impacts of different cluster plot designs on inventory efficiency under a variety of simulated forest heterogeneity scenarios. The main goals of this study were to uncover and interpret the effects of different types and scales of heterogeneity on the relationship between variance and plot design choices, and to provide a conceptual and experimental framework for investigating inventory plot design optimization.

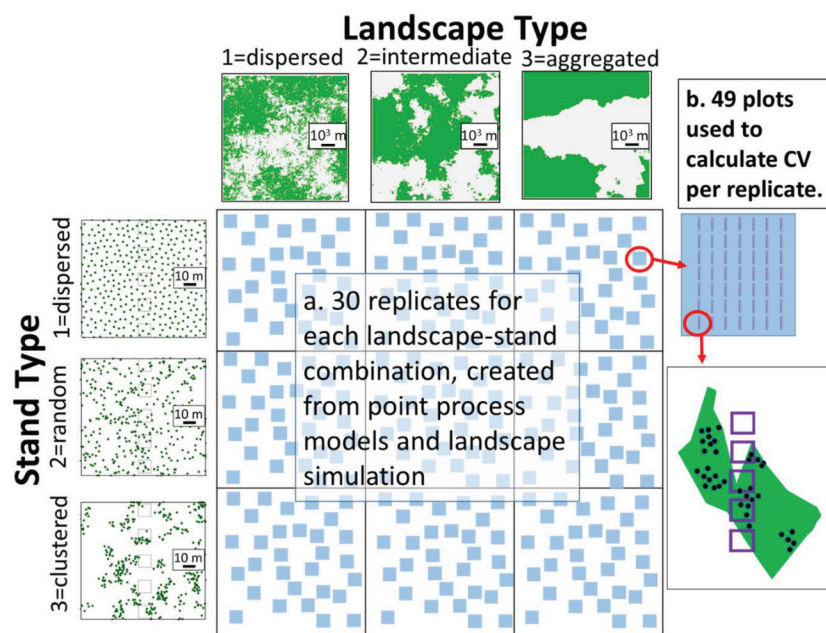
Methods

Simulation experiment

A repeated measures factorial simulation experiment with multiple crossed factors (Oehlert 2000, p. 438) was conducted to model the CV of forest tree density as a function of two spatial pattern factors and three cluster plot design factors. Each of the two spatial pattern factors, representing landscape- (L) and stand- (S) scale heterogeneity, had three levels (Fig. 1):

¹Supplementary data are available with the article at <https://doi.org/10.1139/cjfr-2020-0508>.

Fig. 1. Conceptual model depicting factorial experimental design. (a) Thirty replicates for each of three levels of landscape-scale (L) heterogeneity were simulated for each of three levels of stand-scale (S) heterogeneity. (b) At each plot located on a 7×7 grid superimposed on each landscape, stand-scale patterns were simulated for each level of S . Candidate cluster plot designs were superimposed over each plot location for each replicate, and TPH was calculated per plot. Coefficients of variation (CVs) from the 49 plots were calculated for each combination of candidate plot design and replicate. [Colour online.]



L1: highly dispersed patterns of forest patches with many small, isolated fragments

L2: intermediate levels of aggregation

L3: highly aggregated patterns, with large, continuous patches

S1: highly dispersed (uniformly distributed) pattern of tree locations

S2: completely random pattern

S3: highly aggregated pattern of tree locations, with trees occurring in clumps

The concepts of aggregation, dispersion, and randomness of patterns are dependent upon the scale of analysis and the definitions of patch. For the purposes of this study, patches are defined as geographic areas that are internally homogeneous and possess clear boundaries. In the context of landscapes, aggregation is thus defined as a patch configuration in which patch sizes are large and edge density is small, whereas dispersion is defined as one with smaller patch sizes and larger edge densities (Fig. 1). In the context of point patterns such as the spatial distribution of trees, patch boundaries are much harder to define. We therefore draw on definitions of dispersion and aggregation from the field of point process statistics (Baddeley et al. 2015), in which the patterns are characterized as having distributions of interpoint distances that reflect spatial grouping, randomness, or uniformity.

For each of the nine factor level combinations, 30 replicates were generated via simulation for a total of 270 replicates (procedure described in the section on simulation details and Fig. 1, below).

Each of the 270 replicate heterogeneity scenarios was sampled once at 49 plot locations with a cluster plot design, consisting of a linear array of square subplots aligned north to south (Fig. 1). Each cluster plot design was drawn from a set of designs consisting of every combination of the following factors:

Distance between subplots (d): 10, 25, or 50 m subplot separation (measured between subplot edges)

Number of subplots (m): 2, 3, 4, or 5 subplots per plot

Subplot area (a): 0.01 or 0.2 ha subplot area.

In other words, each member of a set of $3(d) \times 4(m) \times 2(a) = 24$ cluster plot designs was used to sample each of the 270 replicates at the 49 plot locations, as shown conceptually in Fig. 1. This process required repeated measurements, as the 49 plot locations were always the same on each replicate; the only thing that changed was the plot design. The sampling intensity within each replicate was approximately 214 ha per plot, which is within a range of intensities employed by many countries' national forest inventories (Lawrence et al. 2009, p. 41).

The variable recorded for each of the $30 \times 3(L) \times 3(S) \times 3(d) \times 4(m) \times 2(a) = 6480$ observations was CV of forest trees per hectare (hereafter, N), calculated from the 49 plot-level values and using simple random sampling estimators. Cluster plots were treated as a single stage design (Thompson 2012). Our population includes both forest and nonforest areas, which is the same paradigm used by many large area forest inventories, like that of the United States (Bechtold et al. 2005). This implies that all plots are considered to be 100% in the population and accessible, including plots falling partially or completely outside forest patches. CV was chosen as the dependent variable in the model because it is needed to estimate the required sample sizes to achieve allowable error (n_{req} , eq. 1), and is often used to help guide forest inventory design decisions. A summary of this experiment is as follows:

1. Simulate 270 replicates of the forested heterogeneity scenarios (9 L - S combinations $\times$ 30 replicates each, Fig. 1a).
2. Overlay a 7×7 grid of plot locations on each replicate.
3. Sample each of the 270 replicates at the plot locations from step 2 using each of the 24 plot designs, recording CV of N for each case (Fig. 1b), leading to 6480 observations of CV.
4. Build and interpret a linear model of CV as a function of L , S , d , m , and a .

Simulation procedure

Simulation of L

To create a heterogeneity gradient from simulated landscapes, we used the multifractal map generation feature of the qRule landscape analysis software (Gardner 1999, 2017). We created square maps (10 m pixels, 10.24 km sides, 105 km²) with 50% coverage of each of two landcover classes: forest and nonforest. The software generates realistic maps using a fractal algorithm that produces randomized, spatially correlated patterns of land cover, with the option of controlling the level of aggregation. This allows for the creation of each level of *L* described above and in Fig. 1. We calibrated this algorithm such that the three levels of aggregation corresponded to a gradient of approximate forest edge densities of 432.5 m·ha⁻¹ for *L*₁, 55.0 m·ha⁻¹ for *L*₂, and 11.2 m·ha⁻¹ for *L*₃. Forest edge density is defined as the length of the interface between forest and nonforest pixels, divided by the area of the map. For reference, the maximum possible edge density (an alternating forest–nonforest checkerboard pattern of pixels) would be approximately 2000 m·ha⁻¹, and the smallest possible edge density (the perimeter of a square patch of forest that is surrounded by nonforest and occupies half the landscape area) would be approximately 3 m·ha⁻¹. We chose the simulation method and parameters to cover a diverse range of landscape patterns such as those found in northeastern United States temperate forests fragmented by different levels of urbanization and agriculture (e.g., *L*₂ and *L*₃), and those found in agricultural ecosystems like those in Central America with sparse tree clusters of different sizes within natural grasslands (*L*₁); we provide analysis results and examples and code for calculating edge density and forest proportion from existing maps (Supplementary Material S2¹). For each level of *L*, 90 maps (replicates) were simulated, 30 per level of *S*. The raster (Hijmans 2019) and spatstat (Baddeley and Turner 2005) R packages were used to convert the qRule output files to raster objects.

Simulation of S

For the simulations of tree patterns, an *N* of 388 trees·ha⁻¹ was chosen. This is the average tree density of live trees ≥12.7 cm diameter at breast height on forest land for Pennsylvania, according to the USDA Forest Service's Forest Inventory and Analysis (FIA) database (USDA 2020). We considered other commonly reported inventory attributes to use for our study, including basal area per hectare (*G*) and volume per hectare (*V*). *N* was chosen because, in Pennsylvania, the variance of its estimate for trees greater than 12.7 cm diameter at breast height tends to be slightly larger than that for *G* but slightly smaller than that for *V* (USDA 2020). In addition, *N* is a co-equal component of stocking calculations with *G*, and has been considered what is arguably one of the fundamental inventory attributes that FIA produces (Zarnoch and Bechtold 2000). Finally, simulating patterns of *N* is possible using standard point process models that reflect naturally occurring patterns, while simulating *G* and *V* adds complexity by requiring hierarchical models or other techniques that incorporate the effects of tree size in pattern formation.

To create a heterogeneity gradient from three spatial patterns at the stand-scale, spatial point process models were used to simulate the three types of *S* pattern types described above at each plot location using the spatstat R package. Spatial point pattern modelling can be used to both model existing spatial point patterns, and to simulate them, so as to create spatial patterns that might occur in nature (Lister and Leites 2018); Stoyan and Penttinen (2000) describe how various ecological conditions can lead to regular, random, or dispersed patterns of trees. To create the highly dispersed pattern (*S*₁), a simple sequential inhibition pattern generator (rssi) with a 4 m inhibition distance and the requisite number of points to achieve the target *N* was implemented. rssi works by sequentially adding points to the analysis window and rejecting points that fall within the inhibition distance (Baddeley et al. 2015). For the

intermediate level of aggregation (*S*₂), a homogeneous Poisson process pattern generator (rpoispp) was used with the target *N* to create a completely random spatial pattern. rpoispp works by applying a uniform Poisson process within the analysis window to generate complete spatial randomness of points (Baddeley et al. 2015). To create the aggregated spatial pattern (*S*₃), a Thomas cluster process generator (rThomas) was used with a scale parameter of three, a mean number of points per cluster of 10, and the requisite *N* for cluster centers. rThomas works by first generating a uniform Poisson process of initial (parent) points, which are next replaced by clusters of child points that are also generated by a Poisson process, and randomly offset from the parent point location (Baddeley et al. 2015). For each level of *S*, 90 realizations (replicates) were generated, 30 per level of *L*. For each of the 49 plots in each of the *L*–*S* replicates, simulations were performed in a rectangular, 144.7 m × 523.6 m window surrounding each plot center. This window size, which represents a 50 m buffer around the largest candidate cluster plot design's footprint, was chosen to minimize artefacts in the point pattern simulation process that would occur from restricting the simulation to the plot boundaries.

Cluster plot design creation

At each cluster plot location, clusters of different *d*–*m*–*a* configurations at each of 49 locations were superimposed over the simulated *L*–*S* combinations, as shown in Fig. 1b, using the spatstat, raster (Hijmans 2019), and sp (Pebesma and Bivand 2005) packages. Candidate cluster plot designs were located one at a time, and *N* for that cluster plot recorded. Computer code for the R statistical software (R Core Team 2018) for simulating cluster plot designs, landscapes, and stand spatial patterns is provided in Supplementary Material S3¹.

Analysis

To gain insights into how the CV is affected by the different combinations of levels of the variables associated with subplot configuration (*d*, *m*, and *a*) and landscape type (*L* and *S*), we used a mixed effects analysis of variance (ANOVA). This allowed for the testing of main effects and interactions among the variables of interest. In this factorial design, each realization of the simulated landscape and stand pattern combination is a replicate upon which a systematic sample of 49 cluster plots of different design was superimposed to calculate the CV. This required accounting for repeated measures within the factorial design. In addition, the CV was log transformed to minimize heteroskedasticity of residuals. The model form is

$$(2) \quad \log CV_{ijkpqr} = \mu + L_i \times S_j \times m_p \times d_q \times a_r + \gamma_{k(ij)} + \delta_{pk(ij)} + \delta_{qk(ij)} + \delta_{rk(ij)} + \varepsilon_{pqrk(ij)} \sim N(0, \sigma^2)$$

where logCV is the natural log of the coefficient of variation of the estimated *N* of the *k*th replicate, for the *i*th and *j*th levels of *L* and *S*, respectively, and for the *p*th, *q*th, and *r*th level of the plot design variables *m*, *d*, and *a* respectively; μ is the overall mean; γ is the replicate random effect with *k* = 1 to 30 replicates; *L* is the landscape heterogeneity class with *i* = 1 to 3 levels; *S* is the stand heterogeneity class with *j* = 1 to 3 levels; *m* is the number of subplots with *p* = 1 to 4 levels; *d* is the distance between subplots with *q* = 1 to 3 levels; *a* is the subplot area with *r* = 1 to 2 levels; $\delta_{p(k)}$, $\delta_{q(k)}$, and $\delta_{r(k)}$ are random effects accounting for repeated measures for *m*, *d*, and *a*, respectively; and $\varepsilon_{pqrk(ij)}$ is the model error term. The syntax convention we use in eq. 2 was chosen to be consistent with that used elsewhere in the paper, thus parameter names and variable names correspond. After the full model was fit, model terms were evaluated with *F* tests, and those that were not significant at $\alpha = 0.05$ were removed and the model was fit again. A log-likelihood test was performed and results were found to be nonsignificant ($p > 0.05$), suggesting that the reduction of model terms was appropriate (West et al. 2014, p. 220). Version

Fig. 2. Boxplot of the coefficient of variation (CV) values for each landscape ($L1$, $L2$, $L3$) and stand ($S1$, $S2$, $S3$) heterogeneity type. Shaded boxes are interquartile ranges, dark lines are medians, small circles are means, and dashed lines are the range of values for each level. $n = 2160$ for each level. [Colour online.]

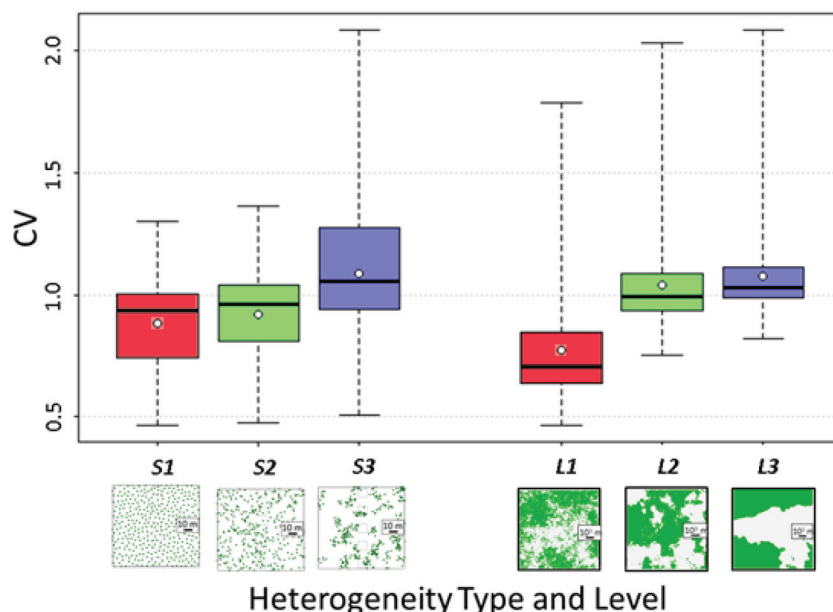


Table 1. Parameters, F statistic, and p values associated with the final model from eq. 2.

Parameter	F value	p value
(Intercept)	66 873.41	<0.0001
a	26 166.63	<0.0001
$a \times S$	5919.29	<0.0001
m	1554.22	<0.0001
$a \times L$	1099.09	<0.0001
L	862.79	<0.0001
d	331.13	<0.0001
$a \times m$	189.4	<0.0001
$a \times m \times S$	144.34	<0.0001
$m \times S$	132.34	<0.0001
$m \times L$	124.42	<0.0001
S	118.66	<0.0001
$d \times L$	45.21	<0.0001
$a \times d$	27.9	<0.0001
$a \times d \times L$	14.4	<0.0001
$a \times L \times S$	14.27	<0.0001
$a \times m \times L$	5.99	<0.0001
$a \times d \times S$	2.95	0.0189
$m \times d$	2.34	0.0293
$d \times S$	1.11	0.3497
$L \times S$	0.42	0.7976

Note: L , landscape type; S , stand type; a , subplot area; m , number of subplots; d , distance between subplots.

1.1-23 of the lme4 package of version 3.5.1 of the R statistical software (Bates et al. 2015; R Core Team 2018) was used to fit the model.

Results

Model results, which take into account the repeated measures design of the experiment and thus allow us to make valid inferences about interactions among design and heterogeneity factors, indicate that heterogeneity type and plot design factors all significantly (F test, $p < 0.05$) affect precision and that many of them interact in different ways (Table 1). In the following sections

we present a summary of the effect of each main factor by averaging across the other factors' levels, and highlight the important interactions found.

Landscape- (L) and stand-level (S) heterogeneity effects on precision

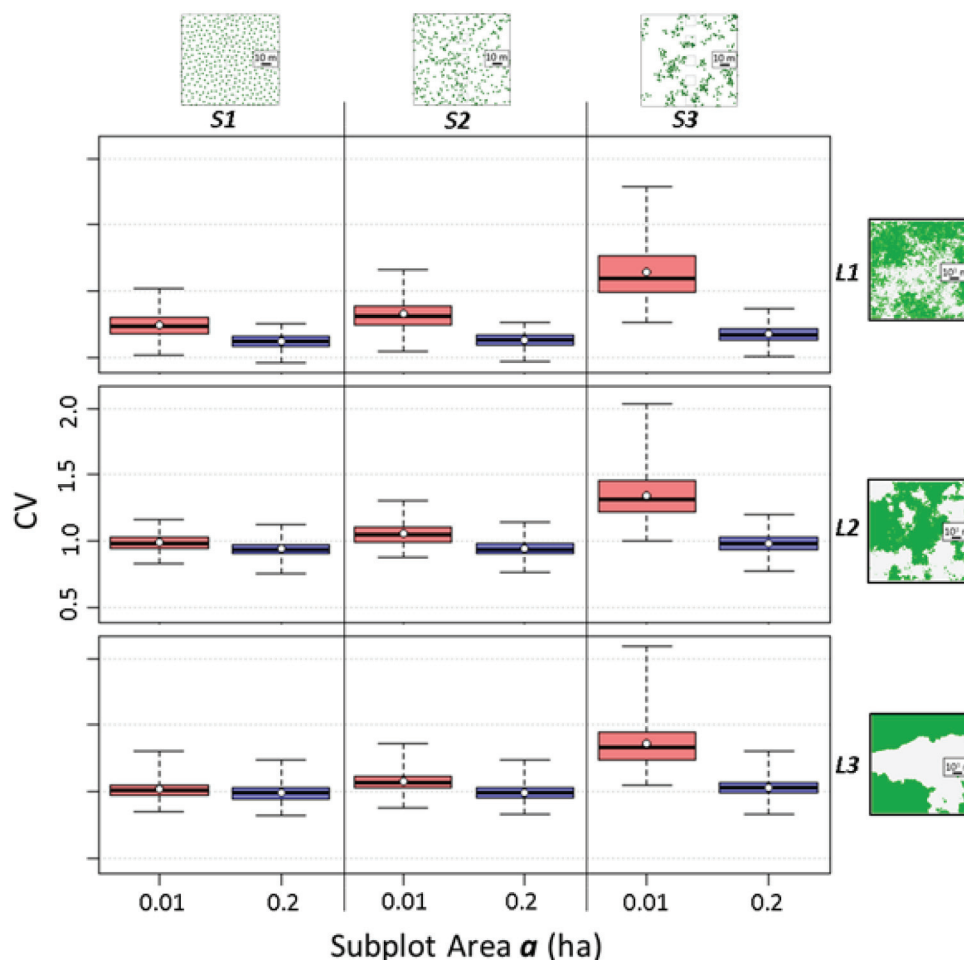
Both landscape- and stand-level heterogeneity had similar effects on precision, with more dispersed patterns ($L1$ and $S1$) having a smaller average CV. The landscape pattern with highly dispersed small patches ($L1$) had a mean CV of 0.77, which is 26% and 28% smaller than those of the more aggregated patches, $L2$ and $L3$, respectively (Fig. 2). Within S , the more dispersed stand pattern ($S1$) had the smallest mean CV, which was 4% and 18% smaller than those for $S2$ and $S3$, respectively (Fig. 2).

Due to the fact that CV is averaged across all levels of the plot design variables, the range of CV variability for each level of S and L is an indicator of the importance that plot design choices can have on CV. At the landscape scale, the CV variability (interquartile range and range) is largest in the more dispersed pattern level ($L1$), and, at the stand scale, largest at the more aggregated level ($S3$, Fig. 2). At the stand scale, the differences in CV variability are greater across levels than at the landscape scale; the range of $S3$ is 90% and 78% larger than those of $S1$ and $S2$, respectively. CV variability across levels of L were all similar (Fig. 2).

Plot design effects on precision

Of the three plot design variables, subplot area had the greatest impact on CV. When subplot area was large, the CV was on average smaller, with the mean CV for $a = 0.01$ equal to 1.06 and that for $a = 0.2$ equal to 0.87 (Fig. 3). In addition, the variability of the CV for the larger subplot was the smallest across most landscape and stand heterogeneity levels, indicating that the two other design variables, m and d , are less influential when subplot area is large. The largest influence of subplot size on the mean CV and variability is in $S3$, the more clustered tree pattern. In contrast, the effect of subplot size is smaller across levels of landscape aggregation, although differences generated by subplot size became smaller from $L1$ (aggregated) to $L3$ (dispersed) (Fig. 3).

Fig. 3. Boxplot representing distribution of coefficient of variation (CV) values for each combination of subplot area (a), landscape type ($L1$, $L2$, $L3$), and stand type ($S1$, $S2$, $S3$). Shaded boxes are interquartile ranges, dark lines are medians, small circles are means, and dashed lines are the range of values for each level. [Colour online.]



Distance between subplots (d) was the least impactful design variable (Fig. 4). What impact it did have generally decreased as landscapes became more aggregated ($L1$ – $L3$) and subplots became large ($a0.01$ – $a0.2$). It was most important in reducing the CV when plot area a was small (0.01 ha) and landscape pattern was highly dispersed and composed of smaller patches ($L1$). In that case, increasing d from 10 to 50 m decreased the CV by over 7% when averaged across levels of S and m . Otherwise, impacts of increasing d had much less or no practical impact on CV compared with changes in the other design factors (Fig. 4).

Number of subplots (m) affected precision by reducing the CV as the number of subplots, and thus total plot size ($m \times a$), increased (Fig. 4). However, the reduction in CV from increasing the number of subplots was more important when subplot area was small and stands had more clustered patterns (Fig. 4). When the subplot area was larger (0.2 ha), the reduction in CV as m increased was of less magnitude and decreased with increasing landscape aggregation from $L1$ – $L3$.

Effects of plot design variables by heterogeneity type on required sample size

Calculating the n_{req} using an AE of 10% and a confidence level of 95% (eq. 1), we present, for each plot design variable, the percentage reduction in n_{req} when increasing the factor level from the lowest to the largest values while averaging across levels of the remaining variables (Fig. 5). Magnitudes of reductions in n_{req}

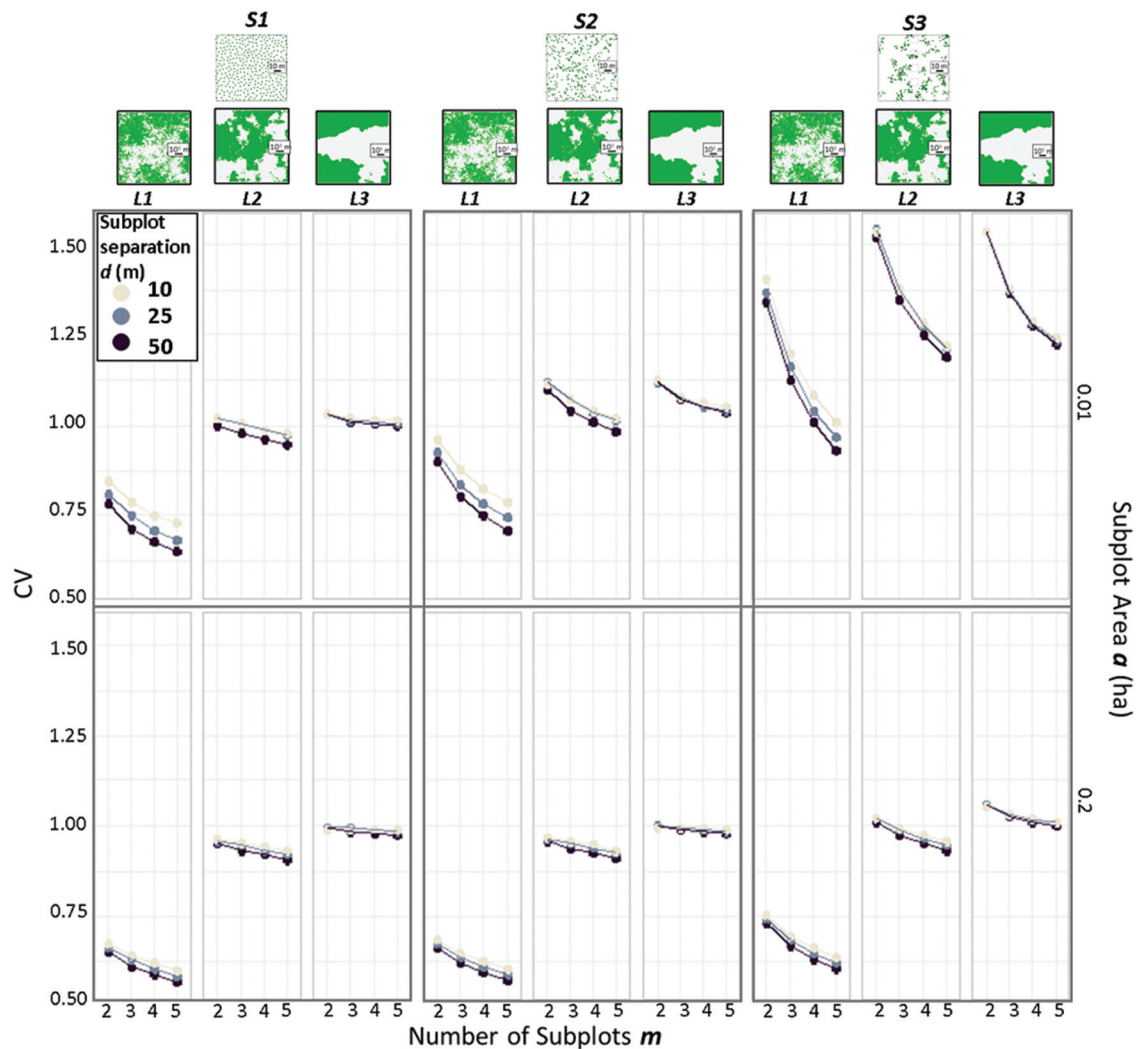
are largest for the most clustered stand patterns ($S3$) and the most dispersed landscape patterns ($L1$), highlighting the importance of plot design choices in those situations. In contrast, smallest reductions were observed for the most uniform stand tree patterns ($S1$) and the aggregated landscape patterns ($L3$). Across heterogeneity scales and levels, the plot design variable that contributed to the greatest reductions on n_{req} was total plot size ($m \times a$) and the least was distance between subplots (d).

Discussion

In this study, we use simulation and a repeated measures factorial experiment to disentangle interactions between variability of an important forest inventory attribute (N), spatial pattern type and scale, and plot design factors. The factor levels of the different plot design variables had important impacts on reductions in the number of sample plots required to achieve AE (Fig. 5). However, these impacts were largely dependent on the landscape L and stand S heterogeneity levels, indicating that the pattern and scale of spatial variability need to be considered when designing plots.

The effects of landscape- and stand-scale heterogeneity on the relationship between plot design and CV are due to the interaction of plot geometry with spatial patterns of tree density. Plots in landscapes with a larger forest edge density ($L1$) are more likely to cross forest patch boundaries and thus contain a mixture of forest and nonforest closer to the average value of the landscape.

Fig. 4. Interaction plots showing the relationship between mean coefficient of variation (CV) and different levels of the factors used in the simulation (*a*, subplot area; *d*, subplot edge separation distance; *m*, number of subplots; *L*, landscape type; *S*, stand type). [Colour online.]



This leads to CVs smaller than those obtained from plots located in aggregated landscape patterns (*L3*), where it is more likely that plots fall entirely either in nonforest areas or forest patches. In *L3*, augmenting either plot area or separation distance therefore does not lead to the acquisition of as much new information as it does in *L1*. This becomes intuitively clear upon inspection of examples of the landscape maps we used in our experiment (Fig. 1).

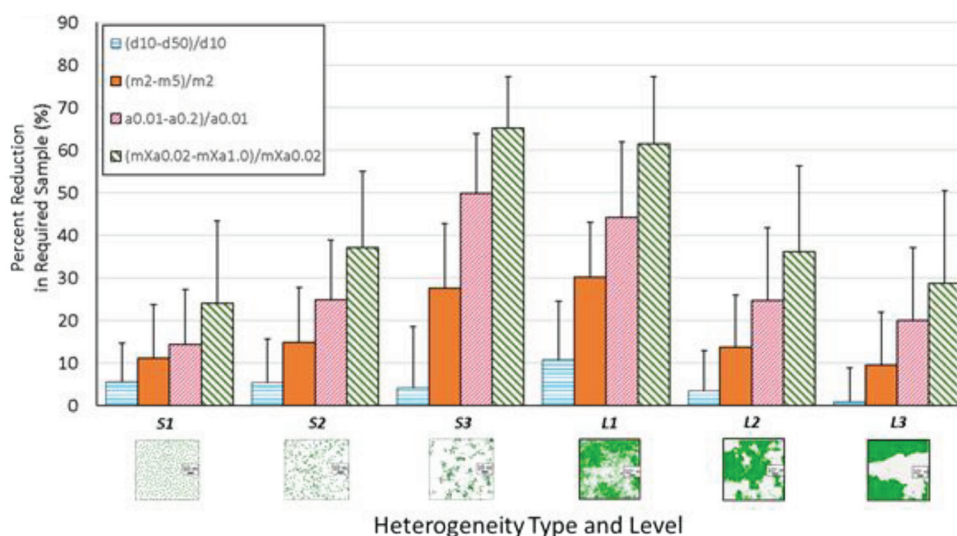
This situation is reversed for levels of stand-scale aggregation. In stands with aggregated patterns (*S3*), plots are more likely to partially fall in either open areas or tree clumps than plots in more uniform stands (*S1* or *S2*). Augmenting plot dimensions or subplot separation in this case has large impacts on CV reduction compared with *S1* or *S2*, where changes in size or subplot spacing will not lead to the plot acquiring new information. This becomes apparent upon inspection of examples of the stand maps shown in Fig. 1; in *S1* and *S2*, tree density is homogeneous at a scale smaller than the dimensions of the subplot. Therefore, there are two spatial pattern-driven mechanisms affecting CV–plot design relationships: one at the landscape scale, where forest patch homogeneity leads to more important plot design effects, and one at the stand scale, where tree pattern aggregation leads to more important effects. We thus conclude that in scenarios like *L1* and *S3*, where increasing the area or subplot spacing of plots is more likely to capture new information, plot design choices have more

impact compared with landscapes like *L3* or *S1*, where the spatial configuration is such that changing plot design parameters within the range we tested is unlikely to dramatically increase the information content of each plot.

Subplot area (*a*) was the most important single plot design factor (Figs. 4 and 5). When averaged across levels of all other factors, the mean CV dropped from 1.06 to 0.87 with larger subplots, which is an 18% decrease. Subplot area also moderates the effects of the other plot design variables, reducing their impacts on CV when the area is larger (Fig. 4). For example, when subplots are small, the increase in *m* leads to a substantial decrease in CV, while when subplots are large, there is a weaker reduction in CV as *m* increases. This was expected, as this aligns with the well-known negative exponential relationship between plot area and relative variance (Smith 1938; Lynch 2017). Our study shows how that relationship changes across a gradient of different types of heterogeneity, becoming more pronounced as stands become more aggregated and less pronounced as landscapes become more aggregated.

Kleinn (1996) found that from a statistical standpoint, subplot separation was a critical factor affecting precision when holding subplot count and area constant; this was due to plots with more internal separation, such as those with subplots configured in a line or L-shape, having smaller intra-cluster correlation. In our

Fig. 5. Percent reduction of required sample size for each plot design variable and heterogeneity type when factor levels increase from lowest to highest and values are averaged across levels of remaining factors. Legend entry indicates how reduction was calculated. Required sample sizes were calculated using eq. 1, with an allowable area (AE) of 10% and a conservative t value of 2.0 for a 95% confidence level, for each heterogeneity type and plot design variable. Error bars represent one standard deviation. [Colour online.]



study, however, subplot separation distance (d) had a relatively small impact on CV and n_{req} in our experiment compared with the other factors (Figs. 4 and 5). The largest impact of d appears for the $L1$ landscape type (across all levels of S), due to the relationship between patch edge density and the set of separation distances we used; $L1$ had a much larger edge density than $L2$ and $L3$. Larger separation distances might have had larger impacts, but these become impractical in the field when using cluster plot designs, and the distances we chose are similar to those employed by other well-established forest inventories like that of the United States (Bechtold et al. 2005). Subplot separation distance d had a larger effect on CV for smaller subplots than for larger subplots, likely because increasing a incorporates so much new information on each plot that the information accrued by increasing d becomes relatively less important.

There are a few practical points to consider. First, although there is a tendency among forest inventory designers and ecologists to prefer larger plots, our results clearly suggest that there is great potential value in choosing subplot area and count based on simulation experiments or pilot studies. For example, Tomppo et al. (2014) used a simulation study that took into account the naturally occurring spatial patterns of existing vegetation maps to find optimal sample and plot design; they repeatedly superimposed cluster plots with different subplot configurations on existing maps in their study area (as opposed to our approach, which involved simulation of a gradient of landscape-scale patterns with known characteristics). Our approach (and our code in the Supplementary Material S2 and S3¹) can easily be adapted to use samples of existing remote sensing-based maps that reflect the patterns in the population under study.

Second, our results suggest that investments in increasing subplot separation distance do not in general have a big impact relative to changing other factors. This supports design choices made by countries that use square, triangular, cross, or L-shaped subplot configurations in their NFIs (Tomppo et al. 2010); there are logistical advantages that can be gained by compact designs in terms of smaller walking distances required to visit all subplots and return to a starting point. However, if using small subplots in landscapes with a large edge density (such as $L1$), subplot separation can lead to meaningful gains in precision by allowing the plot to accumulate new information about the landscape by avoiding

redundant measurements in patches exhibiting autocorrelation. For the same number of subplots, linear arrays allow for the maximum subplot separation.

Finally, it must be noted that logistical or cost considerations associated with a particular design are often very important when making design decisions. For example, integration with remote sensing might be a critical design component if the goal is calibration or validation of map products, or if the data are to be used with some of the more complex designs that benefit from a high degree of spatial alignment between plots and imagery pixels (Næsset et al. 2015); in this case, additional decisions surrounding inferential paradigm and treatment of errors induced by spatial mismatch between subplots and pixels need to be considered.

A second important logistical consideration is the balance between field work costs and those for transportation between plots; when these are considered, relationships between design variables and efficiency will differ (Zeide 1980; Lynch 2017), and efforts should be made to attach cost data to each plot design and model inventory cost as a function of the design and heterogeneity variables, using the same basic methods as those described herein. In the absence of reliable cost information, however, the principles derived from our research will be helpful when making design choices aimed at efficiently achieving a fixed AE.

The conclusions from this study would be difficult to predict without simulation experiments to test plot design options. We chose to use a landscape simulation algorithm that produces landcover patterns that resemble those found on actual landscapes, yet can be parameterized to reflect a heterogeneity gradient (Supplementary Material S2¹; Gardner 1999, 2017). For stand patterns, modelling and stochastic simulation using point process theory allows for the creation of patterns that reflect processes that occur in nature. For example, Lister and Leites (2018) developed a hierarchical point process modelling framework that allows for simulation of realistic patterns of trees that reflect the outcomes of factors like competition and topographic influences. If point process models and landscape maps are calibrated using data from pilot studies within the population of interest, multiple inventory design scenarios can be tested before investing in a specific design. Furthermore, Stoyan and Penttinen (2000) provide a review of pattern types that are associated with various types of forest ecosystem conditions, and these heuristics could aid

practitioners in the selection of tree pattern parameters in new study areas. We have provided R code for plot design and landscape and stand spatial pattern simulation and for subsetting of existing landscape maps (Supplementary Material S2 and S3¹) so that in future work, different plot designs (such as those incorporating remote sensing) and spatial pattern scenarios can be constructed to meet individual needs.

Conclusions

Plot design variables have the largest impact on increasing precision of the estimate when the attribute has a dispersed pattern at the landscape level (L1) and a clustered pattern at the stand level (S3). In these situations, investments in experiments to optimize plot design become highly relevant. If choosing to work with small subplots, then number of subplots and separation distance become more relevant than when considering larger subplots. Large landscape pattern aggregation level leads to a larger CV, regardless of stand pattern heterogeneity and with little effect of plot design variables in changing the overall CV values. In the same way, dispersed patterns at the stand level lead to smaller CVs, and plot design variable choices are of less relevance.

It is difficult to predict the effects of forest heterogeneity on the precision outcomes of different forest inventory plot designs without either using guidelines derived from previous experience and general principles or using a modelling framework that allows for testing hypotheses about the effects of different heterogeneity scenarios on precision. The current study provides both. We have created a simulation framework for creating forest and landscape patterns with different spatial configurations of forest trees at both landscape and stand scales, and we have identified general principles that can be used to guide design choices. When inventory planners confront areas with varying landscape and tree spatial patterns, results from our study will provide heuristics that will help choose from among the wide range of plot design choices, identify which variables to focus on, and gain insight into how spatial pattern changes at two scales can affect design choices. Our goal was to provide monitoring system designers with principles and tools with which to design more efficient inventories.

Acknowledgements

We thank Ephraim Hanks, Eric Zenner, Doug Miller, Charles Scott, James Westfall, Samidha Shetty, Tyler Garner, John Stanovick, and four anonymous reviewers for advice and guidance on earlier versions of this manuscript. This work was partially supported by the USDA National Institute of Food and Agriculture and Hatch Appropriations under Project No. PEN04700 and Accession No. 1019151.

References

- Arvanitis, L.G., and O'Regan, W.G. 1967. Computer simulation and economic efficiency in forest sampling. *Hilgardia*, **38**(2): 133–164. doi:10.3733/hilg.v38n02p133.
- Baddeley, A., and Turner, R. 2005. spatstat: an R package for analyzing spatial point patterns. *Journal of Statistical Software*, **12**(6): 1–42. doi:10.18637/jss.v012.i06.
- Baddeley, A., Rubak, E., and Turner, R. 2015. Spatial point patterns: methodology and applications with R. Chapman & Hall/CRC Interdisciplinary Statistics, New York.
- Bates, D., Mächler, M., Bolker, B., and Walker, S. 2015. Fitting linear mixed-effects models using lme4. *J. Stat. Softw.* **67**(1): 1–48. doi:10.18637/jss.v067.i01.
- Bechtold, W., Scott, C., and Patterson, P. 2005. The forest inventory and analysis plot design. In *The enhanced forest inventory and analysis program – national sampling design and estimation procedures*. Edited by W.A. Bechtold and P.L. Patterson. USDA Forest Service, Asheville, NC. pp. 27–42. Available from http://www.srs.fs.usda.gov/pubs/gtr/gtr_srs080/gtr_srs080-bechtold001.pdf [Accessed 27 November 2020].
- Branthomme, A. 2004. National forest inventory field manual template [working paper]. Forest Resources Assessment Program, Food and Agriculture Organization, United Nations, Rome, Italy. Available from <http://www.fao.org/3/a-ae578e.pdf> [accessed 14 January 2021].
- Brink, G.E., and Schreuder, H.T. 1992. ONEPHASE: a simulation program to compare 1-phase sampling strategies [research paper]. USDA Forest Service, Rocky Mountain Forest and Range Experiment Station, Fort Collins, Colo. Available from <http://catalog.hathitrust.org/Record/011397236> [accessed 27 November 2020].
- Gardner, R.H. 1999. RULE: Map generation and a spatial analysis program. In *Landscape ecological analysis: issues and applications*. Edited by J.M. Klopatek and R.H. Gardner. Springer, New York, NY. pp. 280–303. doi:10.1007/978-1-4612-0529-6_13.
- Gardner, R.H. 2017. Characterizing categorical map patterns using neutral landscape models. In *Learning landscape ecology: a practical guide to concepts and techniques*. Edited by S.E. Gergel and M.G. Turner. Springer, New York, NY. pp. 83–103. doi:10.1007/978-1-4939-6374-4_6.
- Gove, J. 2017. Some refinements on the comparison of areal sampling methods via simulation. *Forests*, **8**(10): 393. doi:10.3390/f8100393.
- Gregoire, T. 1998. Design-based and model-based inference in survey sampling: appreciating the difference. *Can. J. For. Res.* **28**(10): 1429–1447. doi:10.1139/cjfr-28-10-1429.
- Haddad, N.M., Brudvig, L.A., Clobert, J., Davies, K.F., Gonzalez, A., Holt, R.D., et al. 2015. Habitat fragmentation and its lasting impact on Earth's ecosystems. *Sci. Adv.* **1**(2): e1500052. doi:10.1126/sciadv.1500052. PMID:26601154.
- Heilman, G.E., Stritholt, J.R., Slosser, N.C., and Dellasala, D.A. 2002. Forest fragmentation of the conterminous United States: assessing forest intactness through road density and spatial characteristics – forest fragmentation can be measured and monitored in a powerful new way by combining remote sensing, geographic information systems, and analytical software. *Bioscience*, **52**(5): 411–422. doi:10.1641/0006-3568(2002)052[0411:FFOTCU]2.0.CO;2.
- Hijmans, R.J. 2019. raster: geographic data analysis and modeling, version 2.9-23. Available from <https://CRAN.R-project.org/package=raster> [accessed 27 November 2020].
- Hou, Z., Xu, Q., Hartikainen, S., Antilla, P., Packalen, T., Maltamo, M., and Tokola, T. 2015. Impact of plot size and spatial pattern of forest attributes on sampling efficacy. *For. Sci.* **61**(5): 847–860. doi:10.5849/forsci.14-197.
- IPCC. 2006. IPCC guidelines for national greenhouse gas inventories. Institute for Global Environmental Strategies, Japan. Available from <http://www.ipcc-nggip.iges.or.jp/public/2006gl/> [Accessed 27 November 2020].
- Kleinn, C. 1996. Ein Vergleich der Effizienz von verschiedenen Clusterformen in forstlichen Großrauminventuren. *Forstw. Cbl.* **115**: 378–390. doi:10.1007/BF02738616.
- Kleinn, C., Morales, D., and Ramirez, C. 2001. Large area inventory of tree resources outside the forest: What is the problem? In *Proceedings of the IUFRO 4.11 Conference, Greenwich 2001: Forest Biometry, Modelling and Information Science*, 26–29 June 2001.
- Lawrence, M., McRoberts, R.E., Tomppo, E., Gschwantner, T., and Gabler, K. 2009. Comparisons of national forest inventories. In *National forest inventories: pathways for common reporting*. Edited by E. Tomppo, T. Gschwantner, M. Lawrence, and R.E. McRoberts. Springer, New York, NY. pp. 29–42.
- Lister, A.J., and Leites, L.P. 2018. Modeling and simulation of tree spatial patterns in an oak-hickory forest with a modular, hierarchical spatial point process framework. *Ecol. Model.* **378**: 37–45. doi:10.1016/j.ecolmodel.2018.03.012.
- Loetsch, F., and Haller, K.E. 1973. Forest inventory. BLV, Munich, Germany.
- Lynch, A.M. 2003. Comparison of fixed-area plot designs for estimating stand characteristics and western spruce budworm damage in southwestern U.S.A. forests. *Can. J. For. Res.* **33**(7): 1245–1255. doi:10.1139/x03-044.
- Lynch, T.B. 2017. Optimal plot size or point sample factor for a fixed total cost using the Fairfield Smith relation of plot size to variance. *For. Int. J. For. Res.* **90**(2): 211–218. doi:10.1093/forestry/cpw038.
- Mackisack, M.S., and Wood, G.B. 1990. Simulating the forest and the point-sampling process as an aid in designing forest inventories. *For. Ecol. Manage.* **38**(1): 79–103. doi:10.1016/0378-1127(90)90087-R.
- McRoberts, R.E. 2010. Probability- and model-based approaches to inference for proportion forest using satellite imagery as ancillary data. *Remote Sens. Environ.* **114**(5): 1017–1025. doi:10.1016/j.rse.2009.12.013.
- Næsset, E., Bollandsås, O.M., Gobakken, T., Solberg, S., and McRoberts, R.E. 2015. The effects of field plot size on model-assisted estimation of above-ground biomass change using multitemporal interferometric SAR and airborne laser scanning data. *Remote Sens. Environ.* **168**: 252–264. doi:10.1016/j.rse.2015.07.002.
- Oehlert, G.W. 2000. A first course in design and analysis of experiments. W. H. Freeman, New York.
- Pebesma, E.J., and Bivand, R.S. 2005. Classes and methods for spatial data in R. *R News*, **5**(2): 9–13.
- Picard, N., Gamarra, J.G.P., Biragazzi, L., and Branthomme, A. 2018. Plot-level variability in biomass for tropical forest inventory designs. *For. Ecol. Manage.* **430**: 10–20. doi:10.1016/j.foreco.2018.07.052.
- Picard, N., Nouvellet, Y., and Sylla, M.L. 2004. Relationship between plot size and the variance of the density estimator in West African savannas. *Can. J. For. Res.* **34**(10): 2018–2026. doi:10.1139/x04-079.
- R Core Team. 2018. R: a language and environment for statistical computing, version 3.5.1. R Foundation for Statistical Computing, Vienna, Austria. Available from <https://www.R-project.org/> [accessed 27 November 2020].

- Reich, R.M., and Arvanitis, L.G. 1992. Sampling unit, spatial distribution of trees, and precision. *North. J. Appl. For.* **9**(1): 3–6. doi:[10.1093/njaf/9.1.3](https://doi.org/10.1093/njaf/9.1.3).
- Schreuder, H.T., Banyard, S.G., and Brink, G.E. 1987. Comparison of three sampling methods in estimating stand parameters for a tropical forest. *For. Ecol. Manage.* **21**(1): 119–127. doi:[10.1016/0378-1127\(87\)90076-4](https://doi.org/10.1016/0378-1127(87)90076-4).
- Smith, H. 1938. An empirical law describing heterogeneity in the yields of agricultural crops. *J. Agric. Sci.* **28**: 1–23. doi:[10.1017/S0021859600050516](https://doi.org/10.1017/S0021859600050516).
- Stoyan, D., and Penttinen, A. 2000. Recent applications of point process methods in forestry statistics. *Stat. Sci.* **15**(1): 61–78. doi:[10.1214/ss/1009212674](https://doi.org/10.1214/ss/1009212674).
- Thompson, S.K. 2012. *Sampling*. 3rd ed. John Wiley and Sons, Inc., Hoboken, NJ.
- Tomppo, E., Gschwantner, T., Lawrence, M., and McRoberts, R.E. 2010. *National forest inventories: pathways for common reporting*. Springer, New York, NY.
- Tomppo, E., Malimbwi, R., Katila, M., Mäkisara, K., Henttonen, H.M., Chamuya, N., et al. 2014. A sampling design for a large area forest inventory: case Tanzania. *Can. J. For. Res.* **44**(8): 931–948. doi:[10.1139/cjfr-2013-0490](https://doi.org/10.1139/cjfr-2013-0490).
- UNFCCC. 2016. Key decisions relevant for reducing emissions from deforestation and forest degradation in developing countries (REDD+). United Nations, Bonn, Germany. Available from https://unfccc.int/files/land_use_and_climate_change/redd/application/pdf/compilation_redd_decision_booklet_v1.2.pdf [accessed 19 January 2021].
- USDA. 2020. Evaluator retrieval for volume per acre of forest, basal area per acre of forest, and trees per acre of forest for trees greater than 5 inches dbh, Pennsylvania EVALID 422018. Forest Inventory and Analysis, U.S. Department of Agriculture (USDA), Forest Service, Northern Research Station. Available from <http://apps.fs.usda.gov/Evalidator/evalidator.jsp> [accessed 27 November 2020].
- West, B.T., Welch, K.B., and Galecki, A.T. 2014. *Linear mixed models: a practical guide using statistical software*. 2nd ed. CRC Press LLC, Boca Raton, Fla.
- Yim, J.-S., Shin, M.-Y., Son, Y., and Kleinn, C. 2015. Cluster plot optimization for a large area forest resource inventory in Korea. *For. Sci. Technol.* **11**(3): 139–146. doi:[10.1080/21580103.2014.968222](https://doi.org/10.1080/21580103.2014.968222).
- Zarnoch, S.J., and Bechtold, W. 2000. Estimating mapped-plot forest attributes with ratios of means. *Can. J. For. Res.* **30**(5): 688–697. doi:[10.1139/cjfr-30-5-688](https://doi.org/10.1139/cjfr-30-5-688).
- Zeide, B. 1980. Plot size optimization. *For. Sci.* **26**(2): 251–257. doi:[10.1093/forestscience/26.2.251](https://doi.org/10.1093/forestscience/26.2.251).

Copyright of Canadian Journal of Forest Research is the property of Canadian Science Publishing and its content may not be copied or emailed to multiple sites or posted to a listserv without the copyright holder's express written permission. However, users may print, download, or email articles for individual use.