

Simulation räumlicher Verteilungen und Dichteschätzungen von aggregierten Pflanzenpopulationen mit Hilfe von Schätzrahmen (Quadrat sampling)

Kai Schmidt, Nemaplot, Bonn, 2006

Prolog:

Diese Simulationsstudie über die Problematik der Dichteschätzung seltener Pflanzen mit Hilfe von Schätzrahmen ist Bestandteil eines vom Ministeriums für Umwelt und Naturschutz, Landwirtschaft und Verbraucherschutz des Landes Nordrhein-Westfalen gefördertes Forschungsvorhaben **„Erfolgskontrolle des Vertragsnaturschutzes anhand der Populationsgrößen und –entwicklung seltener und gefährdeter Farn- und Blütenpflanzen“**. Dieses Projekt wurde als Kooperationspartner unter der Leitung des Institutes für Nutzpflanzenwissenschaften und Ressourcenschutz (INRES) Ökologie der Kulturlandschaft –Geobotanik und Naturschutz-, Prof. Dr. W. Schumacher im Rahmen des Lehr- und Forschungsschwerpunkt „Umweltverträgliche und standortgerechte Landwirtschaft (USL) der Landwirtschaftlichen Fakultät der Rheinischen Friedrich-Wilhelms-Universität Bonn durchgeführt.

Zitiervorschlag:

SCHUMACHER, W. et al (2006), Erfolgskontrolle des Vertragsnaturschutzes anhand der Populationsgrößen und –entwicklung seltener und gefährdeter Farn- und Blütenpflanzen. Landwirtschaftlichen Fakultät der Rheinischen Universität Bonn, Schriftenreihe des Lehr- und Forschungsschwerpunktes USL, Nr. 148, 160 Seiten.

1 Einleitung

Seltene und gefährdete Farn- und Blütenpflanzen können unter günstigen Voraussetzungen durchaus Populationsdichten von 5000 - 10000 Individuen pro Gebiet erreichen. Diese Abundanz ist in einem vertretbaren Aufwand nur durch entsprechende Stichprobenverfahren zu ermitteln, anhand derer Rückschlüsse auf die Grundgesamtheit der Population gezogen werden können. Unproblematisch ist die Situation bei homogenen Grundgesamtheiten (COCHRAN 1972), was aber selten der Fall ist. Die Individuen einer Population sind nicht gleichmäßig in ihrem Habitat verteilt, sondern neigen zur Klumpungsbildung, die auch als Überdispersion bezeichnet wird. Eine offensichtliche räumliche Heterogenität der hier behandelten Pflanzenarten kann vorausgesetzt werden und muss nicht spezifisch untersucht werden.

Somit sind die Populationskenngrößen durch zwei Parameter charakterisiert: Abundanz und Dispersion. Dichteschätzungen in heterogenen Grundgesamtheiten erfordern besondere Berücksichtigung in der Art der Datenerhebung und Auswertung, um a) eine Güte dieser Schätzung zu definieren, b) Fehlinterpretationen zu vermeiden. Von besonderer Bedeutung ist hierbei, dass bei langfristigen, wiederholten Untersuchungen über die Zeit signifikante Veränderungen in der Dichte der Population erkannt werden, und welcher Aufwand betrieben werden muss, damit diese Änderungen nicht im statistischen Rauschen untergehen.

Ziel dieses Teilprojekts ist es, den Vorgang der Dispersions- und Dichteschätzung von Populationen anhand verschiedener Stichprobenverfahren unter Zuhilfenahme einfacher stochastischer Modelle am Computer zu simulieren, allgemein übliche statistische Auswertungsverfahren auf ihre Anwendungsmöglichkeiten hin zu überprüfen, und eventuelle systematische Fehler zu erkennen. Praktische Grundlage dieser Analyse sind 14 seltene und/oder bedrohte Pflanzenarten aus einer dreijährigen Feldbeobachtung in der Eifel.

Die Vorteile von Simulationen liegen auf der Hand. In kurzer Zeit ist es mit geringem personellen Aufwand möglich, aufwändige Verfahren wie die Erfassung von Populationseigenschaften in großem Wiederholungsumfang durchzuführen, was am Standort unmöglich ist. Bei Anwendung von Simulationsmodellen muss aber auch die Übertragbarkeit auf die Realität berücksichtigt werden, das Ziel sollte sein, realitäts-treue Echtzeitbedingungen anhand von allgemeinen Vorinformationen zu schaffen. Zur Überprüfung der Güte der in der Praxis durchgeführten Stichprobenverfahren wird hier ein rein empirischer Ansatz verfolgt. Über Simulationen werden Populationen verschiedener Dichten und Heterogenität generiert und durch entsprechende Stichprobenverfahren ausgezählt.

Es wird davon ausgegangen, dass diese Analysen nicht zu universell anwendbaren Verfahren führen werden, sondern zu optimierten Stichprobenverfahren unter besonderer Berücksichtigung der physiologischen/phänologischen Besonderheiten einer Art innerhalb ihres Ökosystems.

2 Statistische Grundlagen

Angewandte Verteilungen bei zufälligen und heterogenen Grundgesamtheiten

Um Populationen mit unbekannter räumlicher Verteilung zu charakterisieren, ist man bestrebt eine mathematische Verteilung mit bekannten Eigenschaften an eine aus den Stichproben resultierende, unbekannte Verteilung anzupassen. Wenn die Population zufällig verteilt vorliegt, hat jedes Individuum (hier Pflanze) die gleiche Wahrscheinlichkeit in der Stichprobe gefunden zu werden. Eine adäquate Verteilung zur Beschreibung solcher Populationen ist die Poissonverteilung. Folgt die beobachtete Verteilung in der Stichprobe einer Poissonverteilung, so sind die Individuen einer Population zufällig in der Fläche verteilt (RICHTER et al., 1990).

Da die hier untersuchten Arten i.d.R. stark geklumpt vorliegen, wird nicht näher auf die Poissonverteilung eingegangen.

2.1 Die Negativ-Binomialverteilung (NGB)

Die Negativ-Binomialverteilung, im folgenden als NGB abgekürzt, hat sich als nützliches Modell zur Analyse von Feldversuchen erwiesen, deren Stichprobenvarianzen größer sind als die Mittelwerte. Die Möglichkeiten der Anwendung dieser Verteilung bestehen für weite Bereiche von aggregierten Populationen (SEBER 1986; SOUTHWOOD 1978; TAYLOR et al. 1978).

Die NGB hat folgende Wahrscheinlichkeit x Individuen (Pflanzen) in der Stichprobe zu finden:

$$f(x) = \binom{k+x-1}{x} \cdot \left(\frac{\mu}{k}\right)^x \cdot \left(1 + \frac{\mu}{k}\right)^{-x-k} \quad \text{für } x = 0, 1, 2, \dots \quad \text{Gl. 1}$$

mit dem Erwartungswert und der Varianz: $E_{[x]} = \mu$; $V_{[x]} = \sigma^2 = \mu + \frac{\mu^2}{k}$

Diese Verteilung gehört zu den zweiparametrischen Verteilungen und beschreibt nicht nur zahlreiche biologische Prozesse, sondern kann auch auf vielen Wegen erzeugt werden (BOSWELL & PATIL 1970). Bei Überdispersionen kann k als grober

Maßstab für die Aggregation (Grad der Klumpung) angesehen werden. Damit ergibt sich eine Möglichkeit der Vergleichbarkeit verschiedener Populationen. Mit zunehmender Klumpung der Population wird der Parameter k kleiner. Mit größer werdendem k approximiert die NGB an die Poissonverteilung, beschreibt also eine zufällig verteilte Population. Es sei erwähnt, dass der Parameter k mit der Quadratgröße im Verhältnis zur Biotopegröße variiert, d. h. der Parameter k ist

- nicht generell übertragbar und
- erlaubt keine Rückschlüsse auf eine bestimmte Art von Heterogenität.

Problematisch bei dieser Verteilung ist, dass ein zweiter Parameter aus der Stichprobe geschätzt werden muss. ANSCOMBE (1950) gibt mehrere Verfahren an, um den Parameter k zu schätzen. Die einfachste Methode ist die Schätzung über die Momente. Durch Umwandlung von $V_{[x]}$ erhält man:

$$\begin{aligned}\mu &= \bar{x} \\ k &= \frac{\mu^2}{\sigma^2 - \mu}\end{aligned}\tag{Gl. 2}$$

Es sei erwähnt, dass der Schätzer über die Momente verzerrt (*biased*) ist. Gewöhnlich erfolgt die Parameterschätzung über die Maximum-Likelihood-Methode. Durch Bildung der log-likelihood-Funktion und partieller Ableitung nach k erhält man folgende Gleichung, die iterativ gelöst wird:

$$f(k) = \sum_{x=0}^{x_{\max}} \frac{H(x)}{\bar{x} + k} - n \cdot \log\left(1 + \frac{\bar{x}}{k}\right)\tag{Gl. 3}$$

wobei $H(x)$ die kumulierten Häufigkeiten $> x_i$ sind (RICHTER et al. 1990). Obwohl die Schätzmethode über die Momente verzerrt ist, erreicht es in Abhängigkeit vom Verhältnis von x zu k , also kleine k 's, unter bestimmten Bedingungen eine Genauigkeit von 90 bzw. 98% der Maximum-Likelihood-Methode (ANSCOMBE 1950). Ein Nachteil der ML-Methode, neben der aufwendigen Berechnung, ist die notwendige Belegung von $H_{\max}$. Nur ein umfangreicher Stichprobenumfang erreicht eine ausreichende Klassenbelegung, die wiederum für die Parameterschätzung benötigt werden. Aufgrund der starken Überdispersion der hier vorliegenden Populationen können die Schätzer über die Momente trotz aller Kritik als ausreichend genau angesehen werden.

2.2 Der Vertrauensbereich einer Populationsschätzung

Die Dichteschätzung der hier untersuchten seltenen Arten erfolgt über entsprechende Stichprobenverfahren. Anhand der Stichproben werden Rückschlüsse auf die Ge-

samtpopulation in einem Feld gezogen. Die Frage lautet im Folgenden, wie genau die Population der Größe N durch das entsprechende Stichprobenverfahren geschätzt wird, bzw. welche Streuung des geschätzten Mittelwert um den wahren Mittelwert zu erwarten ist. N wird umso genauer geschätzt, je größer der Stichprobenumfang n ist. Man ist also bestrebt einen Vertrauensbereich für den Populationschätzer anzugeben, was bei zu erwartenden rechtsschiefen Verteilung nicht immer möglich ist. Als mögliche Annäherung zur Beschreibung der Güte einer Schätzung kann der Variationskoeffizient als skalierungsunabhängige Größe verwendet werden. (Richter et al. 1990). Er gibt das Verhältnis der Standardabweichung zum Mittelwert in % an und ermöglicht den Vergleich mehrerer Stichprobenverfahren mit unterschiedlichen Populationsdichten.

2.3 Schätzen des Stichprobenumfangs

A priori ist es unmöglich den Stichprobenumfang zu bestimmen, da das zu erwartende Ergebnis und der Erwartungswert des Fehlers von der Populationsdichte und der Dispersion abhängt. D. h. man ist entweder auf Vorinformationen aus ähnlichen Versuchen angewiesen oder es müssen Voruntersuchungen durchgeführt werden. Falls die Verteilung und ein ungefährender Schätzer für die Dichte bekannt sind, lässt sich der Arbeitsaufwand abschätzen.

SEBER (1982, 1986) bietet anhand der Poissonverteilung und der NGB folgende Möglichkeiten zur a) Bestimmung des Variationskoeffizienten und b) zur Schätzung des Stichprobenumfangs und der Überprüfung der Güte einer Schätzung über den Erwartungswert der Varianz an:

Zufällige (Poisson)-Verteilung

$$V[\hat{N}] = \hat{N} \cdot \sqrt{\frac{S}{n} - 1} \quad \text{Gl. 4}$$

$$C = \frac{1}{\sqrt{\hat{N}}} \cdot \sqrt{\frac{S}{n} - 1} \quad \text{Gl. 6}$$

Negativ - Binomial Verteilung

$$V[\hat{N}] = \hat{N} \cdot \left(\frac{S}{n} + \frac{\hat{N}}{k \cdot n} \right) \quad \text{Gl. 5}$$

$$C = \frac{1}{\sqrt{\hat{N}}} \cdot \sqrt{\frac{S}{n} + \frac{\hat{N}}{k \cdot n}} \quad \text{Gl. 7}$$

Es gelten folgende Bezeichnungen:

$\hat{N}$: Geschätzter Populationsumfang

A: Gesamte Untersuchungsfläche

a: Fläche einer Suchparzelle (Schätzrahmens), in der Regel 1m²

n: Anzahl der Suchparzellen, Stichprobenumfang

S: A/a

C: Variationskoeffizient, Maß für die notwendige Güte der Schätzung

k: Parameter der NGB

SEBER (1982, 1986) verwendet das Verhältnis der Habitatgröße zur Gesamtpopulation. N ist meistens nicht verfügbar, kann aber anhand von Vorinformationen abgeleitet werden.

Löst man Gleichung 6, bzw. 7 nach n auf, lässt sich nach RICHTER et al. (1990) für einen vorgegebenen Variationskoeffizient C, entsprechend einer geforderten Genauigkeit, der notwendige Stichprobenumfang abschätzen. Von der statistischen Seite her ist ein C kleiner 0.1 bzw. 0.05 erstrebenswert.

Tabelle 1: Notwendige Stichprobenumfänge (n) bei gegebener Dichte, räumlicher Verteilung und gewünschter Genauigkeit, Bsp. für N = 12500 und A = 2500, Schätzrahmen = 1 m²;

	Zufällige Verteilung	Geklumpete Verteilung				
		k				
C	n	30	2	1	0.8	0.4
0.2	5	6	18	30	36	68
0.1	20	23	70	120	145	270
0.05	78	94	280	480	580	1080

Welche Güte eine Schätzung haben sollte, hängt vom Nutzen, den Verwendungsmöglichkeiten der Informationen und der gewünschten statistischen Genauigkeit ab, die in der praktischen Anwendung den limitierenden Faktoren (Geld, Zeit, Personal usw.) gegenübergestellt werden müssen.

Der Variationskoeffizient stellt die Streuung in Prozent vom Mittelwert da. Wegen der besseren Vergleichbarkeit der unterschiedlichen Populationsgrößen und Stichprobenverfahren werden die Abweichungen für den geschätzten Mittelwert $\hat{\mu}$ vom tatsächlichen Mittelwert μ ebenfalls in Prozent dargestellt.

Es gilt

$$\mu\% = \frac{\mu - \hat{\mu}}{\mu} \cdot 100 \quad \text{Gl. 8}$$

wobei die erwartungstreuen Schätzer für $\hat{\mu} = \bar{x}$ und für die Varianz $\hat{\sigma}^2 = s^2$ ist. Beide Parameter werden aus der Stichprobe geschätzt.

2.4 Stichprobenverfahren

Die statistischen Kenngrößen sind meistens erst bei sehr großen Stichproben exakt. Wie aus Tabelle 1 zu ersehen ist, sind schon bei leichter Heterogenität im Untersuchungsmaterial große Stichproben notwendig, um einen akzeptablen Vertrauensbereich zu erhalten. Dabei sind in den theoretischen Ableitungen keine Aussagen über die Art der Stichprobenerhebung getroffen worden. Die Frage besteht im Folgenden, ob es möglich ist, durch die räumliche Anordnung eines Stichprobenverfahrens die zu erwartende Varianz einzuhalten, bzw. zu verringern. Die Stichproben werden in der Regel mit einem 1m² großen Schätzrahmen durchgeführt. Eine Auswahl möglicher Stichprobenverfahren wird kurz vorgestellt (BUCKLAND et al. 1993; COCHRAN 1972; KUNO 1976; RICHTER et al., 1990; RIPLEY 1981; SCHEAFFER et al. 1979; SEBER 1986):

Tabelle 2: Überblick über Klassen von Stichprobenverfahren

Zufallsstichproben	Zählrahmen werden per gleichverteilten Zufallsvariablen auf der Fläche verteilt
Systematische Stichproben	Der 1. Punkt wird zufällig gewählt, alle weiteren Quadrate folgen einer systematischen oder regulären Verteilung im Verhältnis zum Ausgangspunkt.
Geschichtete Zufallsstichproben, <i>stratified random sampling</i>	Die Untersuchungsfläche wird in gleichgroße Untereinheiten (Schichten, Blöcke) eingeteilt, gleichverteilte Zufallsstichproben/ Untereinheit werden gezogen.
Klumpungsstichproben	Untersuchungsfläche wird in eine feste Anzahl gleicher Untereinheiten unterteilt. Einzelne Blöcke werden entsprechend der Aggregation einer Population weiter unterteilt und Zufallsstichproben gezogen.
Adaptive Verfahren	Anlage wie <i>stratified sampling</i> , aber die Anzahl der Zufallsstichproben richtet sich anhand des Ergebnisses. Wird ein kritischer Wert überschritten, wird weiter gesammelt.

Transekte, Suchstreifen	Anstelle von Zählrahmen werden Suchpfade als Stichprobe verwendet. Einfach einzurichten, aber die Formulierung der Auffindwahrscheinlichkeit kann zu Schwierigkeiten führen.
Distanz Methoden	Es wird die räumliche Distanz von zwei Individuen oder eines Zufallspunktes zum nächsten oder übernächsten Individuum gemessen und entsprechend auf die Gesamtpopulation hochgerechnet.

Nachtrag zur Schicht- oder Blockanlage

Die Blockbildung bezweckt eine Homogenisierung der Untereinheiten wenn eine heterogene Grundgesamtheit vorliegt. Ist eine Schicht homogen, d. h. unterscheiden sich die Messwerte in den Zählrahmen nur wenig, so kann der Mittelwert jeder Schicht mit einer kleinen Stichprobe exakt geschätzt werden. Diese einzelnen Mittelwertschätzungen können dann zu einem genauen Schätzwert der heterogenen Grundgesamtheit zusammengefasst werden (COCHRAN 1972; BROWN & ROTHERY 1993).

$$\hat{\mu} = \frac{1}{A} \sum_{i=1}^k A_i \bar{x}_i \quad \text{Gl. 9}$$

Der Stichprobenmittelwert $\bar{x}_i$ des i-ten Blocks A_i (Stratum) ist ein unverzerrter Schätzer der Block-Grundgesamtheit, daher wird der Mittelwert der Population durch die gewichteten Blockmittel mal der Blockgröße A_i geschätzt. A ist die Summe aller A_i 's.

3 Methoden

3.1 Entwicklung des Simulationsprogramms

Zur Lösung der zu behandelnden Fragestellungen ist ein Computerprogramm entwickelt worden, welches

- die gewünschten Dichten und räumlichen Verteilungen einer Pflanzenpopulation auf einem gegebenen Standort simuliert;
- entsprechend der oben beschriebenen Stichprobenverfahren die Zählrahmen verteilt.;

- Nach Bedarf die entsprechende Blockbildung der Untersuchungsfläche ermöglicht;
- das Auszählen mit Hilfe dieser Zählrahmen simuliert;
- die statistischen Kenngrößen und Verteilungen für eine Simulation berechnet;
- die entsprechenden Verfahren beliebig oft wiederholt und die benötigten Statistiken über die Wiederholungen erzeugt;

3.2 Simulation räumlicher Verteilungen

Zur Erzeugung heterogen verteilter Population wird ein modifizierter Thomas (MT-) Prozess verwendet (Diggle, 1983). In der Theorie wird von einem zufällig verteilten Elternpunkt ausgehend eine normalverteilte Anzahl von Nachkommen erzeugt. Die Distanz R der Nachkommen zum Elternpunkt ist wiederum normalverteilt mit Mittelwert 0 und gegebener Varianz. Je kleiner die Varianz, desto geklumpfter ist die Verteilung. Ein hoher Varianzwert generiert eine mehr zufällige Verteilung um den Elternpunkt. Die endgültige Position eines Nachkommens ergibt sich aus einem gleichverteilten Zufallswinkel ϕ im Intervall von 0 bis 360° . Das Verfahren wird solange wiederholt bis die gewünschte Populationsgröße N erreicht ist. Abb. 1 verdeutlicht die Prozedur.

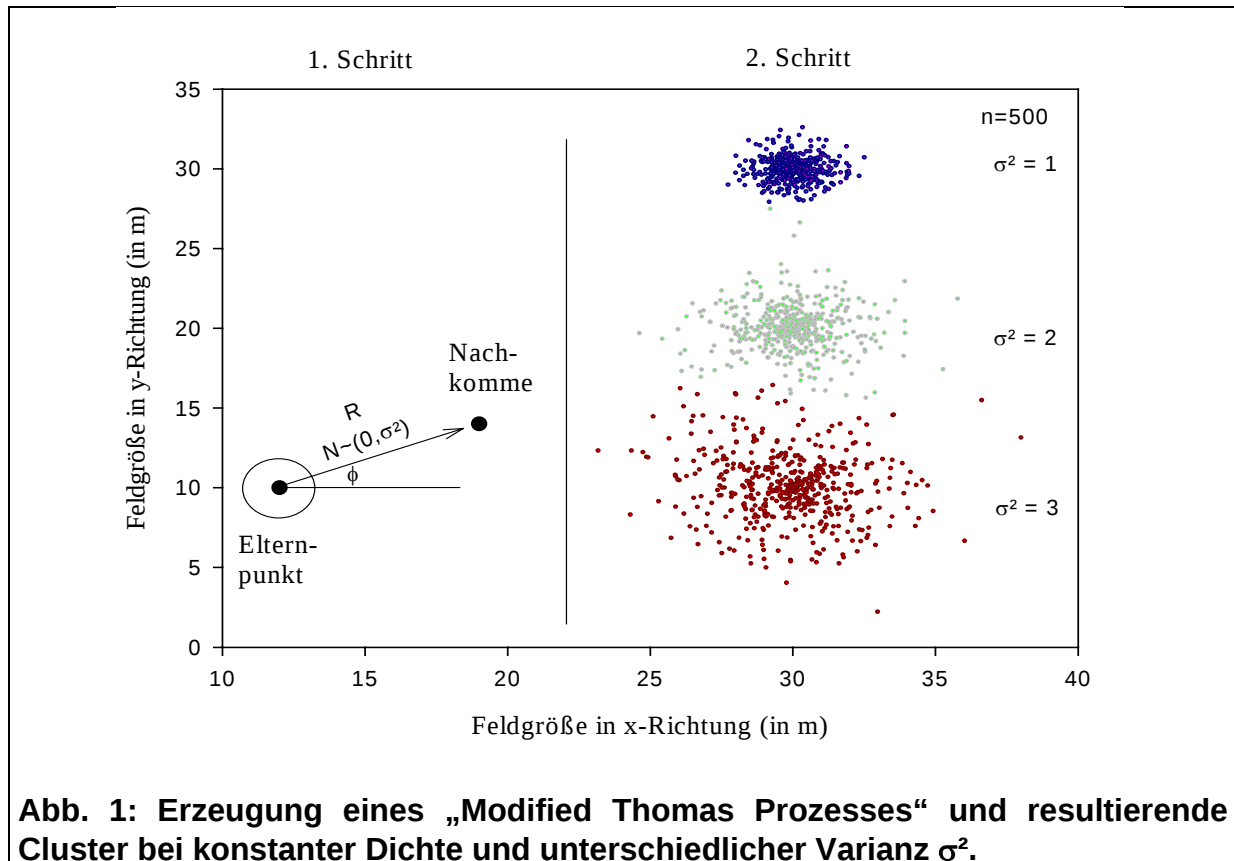


Abb. 1: Erzeugung eines „Modified Thomas Prozesses“ und resultierende Cluster bei konstanter Dichte und unterschiedlicher Varianz σ^2 .

Um möglichst realitätsnahe Verteilungen zu erzeugen, wurden die Populationen anhand der räumlichen Informationen aus den parallel durchgeführten Praxisuntersuchungen erzeugt. Da sowohl die Georeferenz der Feldgröße als auch Lage der Untersuchungsquadrate dokumentiert wurden, lassen sich die Elternpunkte relativ zur Feldrandlage positionieren. Die Anzahl der Nachkommen orientiert sich anhand der vorgefunden Anzahl von Pflanzen in den Quadraten.

Es wurden aus den praktischen Untersuchungen insgesamt 14 repräsentative Kombinationen von Standorten und Pflanzenarten ausgewählt. Entsprechend der Ergebnisse aus den georeferenzierten Informationen wurden mit Hilfe des oben beschriebenen MT-Prozeß Populationen unterschiedlicher Dichte und räumlicher Verteilung generiert, die soweit wie möglich den Erfahrungen aus der Praxis entsprachen. Um eine mögliche Verzerrung der Ergebnisse durch das Verhältnis von Schlaggröße zu Zählrahmengröße zu vermeiden, werden die Zählrahmen im Verhältnis zur Feldgröße dargestellt und ausgezählt. Des Weiteren werden die Schläge auf ein Rechteck idealisiert. In der Realität entsprechen die Untersuchungsflächen in keinem Fall einer

solchen Idealisierung. Eine Übersicht der simulierten Populationen an den verschiedenen Standorte ist in [Anlage A](#) zusammengestellt.

Die beschriebenen Komponenten ergeben damit die Variablen, anhand derer die Frage der Qualität/Optimierung des Stichprobenverfahrens, der -verteilung und des – umfangs zu klären sind, um die eigentliche Zielgröße „Erkennen der Änderung in der Bestandesdichte“ zu bestimmen.

3.3 Stichprobenverfahren

Aus den oben beschriebenen Klassen von möglichen Verfahren wurde folgende Liste zusammengestellt (Tab. 3). Die Wahl der Verfahren orientiert sich anhand der in der Praxis weit verbreiteten (zufällig, Block) und den in diesem Projekt durchgeführten Typen der praktischen Felderhebungen (systematische Verfahren). Auf der Suche nach Verfahren mit reduziertem Aufwand werden innerhalb der systematischen Verfahren mehr empirische Ansätze untersucht (Diagonale Verteilung). Letztendlich wurden noch die adaptiven Zwei-Schritt-Verfahren in die Untersuchungsliste mit aufgenommen.

Tabelle. 3: Übersicht und Aufbau der simulierten Stichprobenverfahren

Nr.	Verfahrensklasse	Verfahren	Beschreibung
1	Zufällig	Zufällig	Quadrate werden zufällig in der Fläche durch gleichverteilte Zufallszahlen verteilt, Überschneidung werden verworfen und neu ausgeteilt.
2	Zufällig	Mindestabstand	Verteilung durch gleichverteilte Zufallszahlen, Quadrate, die einen vorgegebenen Mindestabstand unterschreiten, werden verworfen und neu ausgeteilt.
3	Zufällig	Feste Anteile	Wie zufällig, aber ein fester Anteil des Stichprobenumfangs wird subjektiv in Bereichen mit hoher Dichte gelegt und ausgezählt.
4	Systematisch	Äquidistante Abstände	Die Anzahl der Quadrate werden proportional auf der Gesamtfläche verteilt
5	Systematisch	Schrittweise Abstände im	Erstes Quadrat wird zufällig im Bereich unten/links gelegt, Abstand der folgenden Quad-

		Hoch-Rechts Verfahren	rate wird durch normalverteilte Zufallszahlen erzeugt, Lage im Hoch-Rechts Verfahren. Ist die obere Feldbegrenzung erreicht, wiederholt sich der Prozess rückwärts mit entsprechend korrigierten Abständen bis die gewünschte Anzahl der Quadrate erreicht ist.
6	Systematisch	Felddiagonale	Zuerst wird die Felddiagonale belegt, dann die obere oder untere Diagonale bis die Anzahl der Quadrate erreicht ist.
7	Systematisch	„Fischgräte“	Das Feld wird zufällig in 2 (oder mehr) Bereiche eingeteilt. Entlang der Unterteilung wird zufällig in rechts oder links unterschieden und Quadrate werden bis zur Feldgrenze gelegt. Wiederholt sich bis vorgegebene Anzahl der Quadrate erreicht ist. Abstand der Quadrate zueinander durch normalverteilte Zufallszahlen.
8	Schichtung/Block	Große Blöcke	Die Gesamtfläche wird in große, wenige Blöcke eingeteilt. Innerhalb der Blöcke werden zufällig eine vorbestimmte Anzahl von Quadrate verteilt. Auswertung erfolgt durch entsprechende Blockgewichtung.
9	Schichtung/Block	Kleine Blöcke	Die Gesamtfläche wird in kleinere, viele Blöcke eingeteilt. Innerhalb der Blöcke werden zufällig eine vorbestimmte Anzahl von Quadrate verteilt. Auswertung erfolgt durch entsprechende Blockgewichtung.
10	Schichtung/Block	Variable Schichtgröße	Innerhalb der Fläche werden Blöcke entsprechend von Vorinformation gelegt, wobei auf eine homogene Verteilung innerhalb des Blockes geachtet wird. Auswertung erfolgt durch entsprechende Blockgewichtung.
11	Adaptive, Zweischritt Verfahren	Große Blöcke, kleine Blöcke mit sequentieller Kombination.	Die Gesamtfläche wird in Blöcke eingeteilt. Innerhalb der Blöcke werden zufällig eine vorbestimmte Anzahl von Quadrate verteilt. Die Quadrate werden ausgewertet und je nach Dichte werden weitere Quadrate gelegt bis

		Erläuterung siehe Abb. 2	entweder eine kritische Dichte unterschritten ist oder die relative Änderung einen vorgegebenen Grenzwert unterschreitet. (Sequentielles Verfahren). Auswertung erfolgt durch entsprechende Blockgewichtung.
12	Suchpfad	Strip-Transect	Ausgehend von einem Zufallspunkt am Feldrand wird ein Suchpfad gelegt bis vorbestimmte Länge erreicht ist. Inhalt des Suchpfads wird gezählt und in Beziehung zur Länge des Versuchspfads gesetzt.
13	Suchpfad		Quadrate werden entsprechend des Versuchspfads gelegt. Abstand der Quadrate durch normalverteilte Zufallszahlen.

Erläuterung zu den Adaptiven oder Zweischritt Verfahren

Auch hier wird die Gesamtfläche in Blöcke aufgeteilt. Innerhalb dieser Blöcke wird zufällig eine Stichprobe gezogen. Wenn das Mittel dieser Stichprobe einen zu definierenden Grenzwert (z.B. $\mu > 5$) überschreitet, werden adaptiv weitere Stichproben gezogen. Die Frage ist, wie oft weitere Quadrate ausgezählt werden müssen. Dazu wurde ein sequentielles Verfahren entwickelt und empirische Abbruchkriterien definiert.

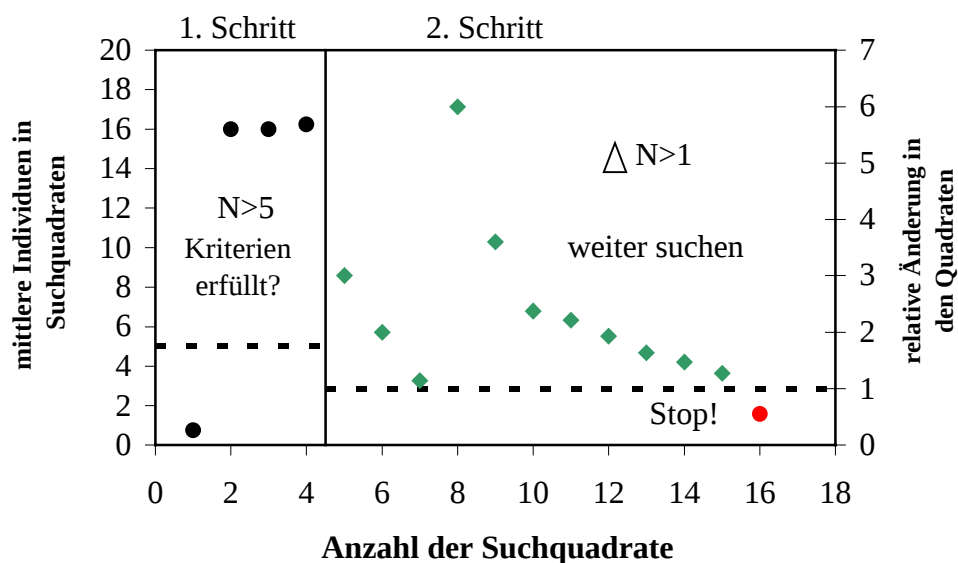


Abb. 2: Sequentieller Aufbau innerhalb eines Zweischritt Stichprobenverfahren

In Abb. 2 sind die Kriterien für ein sequentielles Verfahren dargestellt. Pro Block sind vier Quadrate vorgesehen. Der Mittelwert überschreitet eine vorgegebene kritische Grenze und es folgt der zweite Schritt. Es wird ein weiteres Quadrat zufällig im Block gelegt und ausgezählt. Die Problematik besteht darin, dass es sich erwartungsgemäß um hohe Besatzdichten handelt, deren Auszählung sehr Zeit- und Arbeitsaufwändig sind. Daher gilt es Abbruchkriterien zu definieren. Entscheidungskriterium ist, wann die relative Änderung des Mittelwertes von μ_{n-1} zu $\mu_n < 1$ ist. n ist der Stichprobenumfang. Im obigen Beispiel sind die Abbruchkriterien bei der Stichprobe 7 nahezu erfüllt. Im 8. Quadrat werden sehr viele Individuen gefunden und die relative Änderung überschreitet wieder die Grenze. Aufgrund dieses Anstiegs sind 8 weitere Stichproben notwendig, ehe die Grenze letztendlich unterschritten wird. Der Vorteil besteht darin, dass eine Folge von voll besetzten Quadraten ebenso wie leere Quadrate die Abbruchkriterien der Regel schneller erfüllt als in dem Beispiel dargestellt.

3.4 Überprüfung der simulierten Erwartungswerte

Die Genauigkeit der Simulationen wird anhand einer zufällig verteilten Population getestet, um a) die korrekte Auswertung und Programmierung des Stichprobenverfahrens innerhalb des Simulationsprogramms zu testen, und b) die beobachtete Streuung der Variationskoeffizienten mit der theoretisch zu erwarteten Streuung zu vergleichen.

Tabelle 4: Erwartungswerte bei einer zufällig verteilten Population und zufällige Verteilung der Suchquadrate

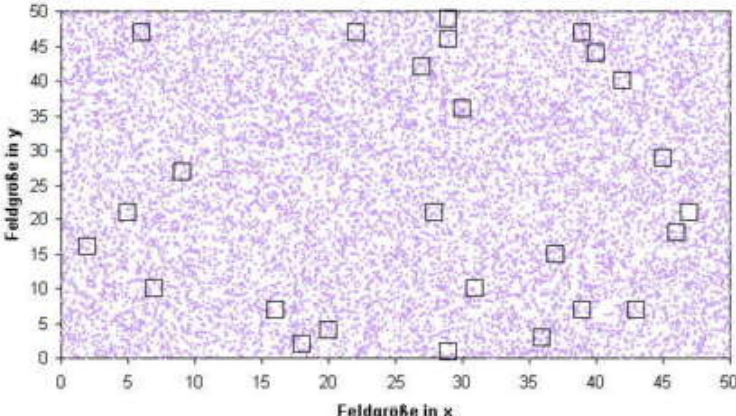
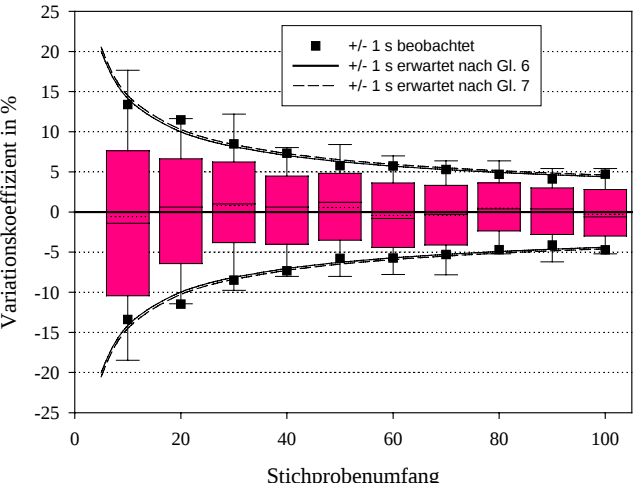
	Population generiert durch gleichverteilte Zufallszahlen, Beispiel: Zufälliges Stichprobenverfahren
	Auswertung durch zufällige Verteilung der Quadrate Erläuterung: $s = \pm 1$ Standardabweichung Streuung beschreibt den Standardfehler des Mittelwertes $s_{\bar{x}}$ Boxplots repräsentieren die 25/75 Quartile, Fehlerbalken die 5/95 Quartile Zum Vergleich: Erwartete Streuung der Mittelwerte nach Gleichung 6, bzw. 7 und beobachtete Streuung aus 200 Wiederholungen. Die 5/95% Quartile präsentieren die Vergleichsgrößen der statistischen Tests.

Abb. 4: Simulierter und erwarteter Standardfehler bei zufälliger Verteilung

Die Darstellung über den Median liefert aufgrund der hier vorliegenden linksgipfeli- gen Verteilungen eine realistischere Abbildung der Analysesituation als es der Mit- telwert könnte. In Abbildung 4 sind sowohl der Mittelwert und der Median aus 200 Wiederholungen dargestellt. Beide streuen marginal um den transformierten Popula- tionsschätzer (Gl. 8), d. h. der Erwartungswert wird eingehalten, falls eine zufällige Population gezählt wurde. Das Beispiel in Abbildung 4 demonstriert weiterhin, welche

Güte die zu untersuchenden Stichprobenverfahren einhalten sollen, falls es sich um eine heterogene Population handelt.

3.5 Versuchsaufbau

Um die Potentiale und Möglichkeiten dieser Verfahren zu testen und untereinander zu vergleichen, wurden für jede der 14 vorgestellten Populationen 10 Stichprobenverfahren bei einem Stichprobenumfang von 10 bis 100 Quadraten mit einer Schrittweite von 10 Quadraten pro Untersuchungsfläche getestet. Jede Testreihe wurde 200-mal wiederholt, die Abweichung vom theoretisch zu erwartenden Mittelwert berechnet und graphisch in den schon beschriebenen Box-plots dargestellt.

4 Ergebnisse

4.1 Simulationen der räumlichen Verteilungen

In Anlage A sind die erzeugten Verteilungsmuster für jeder der hier untersuchten Arten und Standorte zusammengestellt. Dem Erzeugungsprozeß liegen zwar sehr einfache Strukturen zugrunde, aber mit wenigen Parametern sind vielfältige Verteilungsmuster möglich. Das Spektrum reicht von angenähert zufälligen Verteilungen bis hin zur wesentlich größeren Gruppe der extrem starken Clusterung, und kann als gute Annäherung an die Realität angesehen werden. Die Anpassungen an die NGB stellen über den Klumpungsparameter k das Ergebnis numerisch da. In fast allen Fällen zeigt sich die NGB als ausgezeichnetes Modell, um die Wahrscheinlichkeit für die Anzahl der Individuen in einem Zählrahmen vorherzusagen (Anlage A). Daher wird vermutet, dass der Parameter k relativ gut den Dispersionsgrad widerspiegelt. Wie zu ersehen ist, liegt k mehrheitlich im Bereich 0.4 bis 0.7, Extreme gehen runter bis auf 0.1. Selbst wenn man starke Abstriche in der geforderten Genauigkeit macht, kann man anhand dieser Werte abschätzen, wie extrem umfangreich der Stichprobenumfang sein müsste (vgl. Tab. 1), um einen akzeptablen Vertrauensbereich für einen Populationsschätzer zu erhalten.

4.2 Simulationsergebnisse der Stichprobenverfahren

Die Mittelwerte der Populationsschätzer aus 200 Wiederholungen sind den theoretischen Erwartungswerten nach SEBER (1982) gegenübergestellt worden. Die zu erwartende Streuung des Standardfehlers bietet sich als Vergleichsgröße an, um die einzelnen hier durchgeführten Verfahren insgesamt und im Verhältnis zueinander zu testen. Für alle Verfahren gilt erwartungsgemäß, je größer der Stichprobenumfang desto kleiner der Standardfehler, dargestellt in den schon beschriebenen Boxplots.

Um insbesondere zwischen den Verfahren zu vergleichen, werden die Simulationsergebnisse bonitiert. Die Bewertungskriterien richten sich nach der Größe der Streuung, dem Bias oder Abweichung vom Erwartungswertes und dem allgemeinen Verhalten im Verhältnis zum Erwartungswertes der NGB mit berechnetem Klumpungsparameter k . Eine visuelle Analyse aller Boxplots zeigt deutlich, dass es unter den stark geklumpten Populationen mit Stichproben <30 zu großen systematischen Abweichungen kommt. Daher werden die Bonituren nur anhand der Ergebnisse im Bereich der größeren Stichprobenumfänge durchgeführt. Die Boniturstufen sind in Tabelle 5a dargestellt, die Boniturergebnisse in Tabelle 5.

Tabelle 5a

Boniturstufen	Quantile					
Streuung der Quantile im Vergleich zum Erwartungswert der NGB	>>50%	>50%	<50%	<67%	<95%	>95%
%Bias bei mittleren/hohen Stichproben	>30%	>20%	>10%	>5%	0%	
Allgemeines Verhalten / Stichprobenumfang	0	1	2	3	4	
Punkte	0	1	2	3	4	5

Tabelle 5: Boniturübersicht nach Verfahren und Art/Standort

Verfahren	Art	Deutscher Name	<i>Botrychium lunaria</i>	<i>Euphrasia frigida</i>	<i>Gentiana pneumonanthe</i>	<i>Gentiana germanica</i>	<i>Narthecium ossifragum</i>	<i>Narthecium ossifragum</i>	<i>Ophioglossum vulgatum</i>	<i>Orchis morio</i>	<i>Pedicularis sylvatica</i>	<i>Pedicularis sylvatica</i>	<i>Pulsatilla vulgaris</i>	<i>Pulsatilla vulgaris</i>	<i>Narcissus pseudonarcissus</i>	<i>Platanthera chlorantha</i>	Summe:	Rang: Beurteilung der Verfahren
			Mondraute	Nordischer Augentrost	Lungen-Enzian	Deutscher Enzian	Moorlilie	Moorlilie	Natternzunge	Kleines Knabenkraut	Wald-Läusekraut	Wald-Läusekraut	Küchenschelle	Küchenschelle	Gelbe Narzisse	Weißer Waldhyazinthe		
Zufällig	Varianz		3	3	1	3	5	5	3	4	3	1	3	5	3	2	137	5
	Bias		2	4	2	3	4	4	3	4	4	1	4	3	4	2		
	Verhalten		3	4	2	3	4	4	4	4	4	2	4	4	4	3		
Mindest Abstand	Bonitur		8	11	5	9	13	13	10	12	11	4	11	12	11	7	113	8
	Varianz		3	3	1	3	5	5	5	4	3	1	4	5	3	3		
	Bias		1	3	1	1	3	2	2	2	3	0	3	3	2	1		
Aufteilung auf Gesamtfläche	Verhalten		2	4	2	2	4	3	3	3	3	0	4	4	2	2	114	7
	Bonitur		6	10	4	6	12	10	10	9	9	1	11	12	7	6		
	Varianz		3	3	1	3	5	5	5	3	4	1	4	5	3	3		
Hoch-Rechts Schritt	Bias		1	3	1	1	3	2	3	2	3	1	3	2	2	2	130	6
	Verhalten		2	4	2	2	4	3	3	2	4	1	3	3	2	2		
	Bonitur		8	12	6	11	13	10	11	12	3	2	9	13	9	11		
Diagonal	Varianz		3	5	3	5	0	3	0	0	3	3	5	2	5	1	60	10
	Bias		0	3	0	4	0	0	0	0	0	0	3	0	0	0		
	Verhalten		0	3	0	3	0	1	0	0	0	0	4	0	1	0		
Fischgräte	Bonitur		3	11	3	12	0	4	0	0	3	3	12	2	6	1	66	9
	Varianz		1	3	0	0	0	4	2	2	2	0	3	3	2	2		
	Bias		1	2	1	1	1	3	3	0	1	0	2	1	4	0		
Große Blöcke, zufällig im Block	Verhalten		2	3	1	0	0	4	2	1	2	0	2	1	3	1	138	3
	Bonitur		4	8	2	1	1	11	7	3	5	0	7	5	9	3		
	Varianz		3	4	1	3	4	5	4	4	2	1	4	4	3	4		
Kleine Blöcke, zufällig im Block	Bias		2	4	2	3	3	4	4	4	2	2	4	4	3	4	148	1
	Verhalten		3	4	3	4	4	4	2	4	3	3	4	4	3	3		
	Bonitur		8	12	8	11	13	12	8	12	11	8	12	13	10	10		
Adaptiv, große Blöcke	Varianz		3	4	2	3	3	5	4	4	2	2	4	4	4	4	138	3
	Bias		2	4	3	3	3	4	4	4	2	2	3	4	4	3		
	Verhalten		2	4	3	4	3	3	4	4	2	2	3	4	4	3		
Adaptiv, kleine Blöcke	Bonitur		7	12	8	10	9	12	12	12	6	6	10	12	12	10	139	2
	Varianz		3	4	4	4	5	5	5	4	3	2	4	4	3	4		
	Bias		2	4	3	3	4	4	4	4	1	2	4	4	4	2		
Resümee der Standorte/ Arten	Verhalten		2	4	3	3	3	3	4	3	2	1	3	3	3	3		
	Bonitur		7	12	10	10	12	12	13	11	6	5	11	11	10	9		
Resümee der Standorte/ Arten	Summe:		65	110	56	85	95	107	94	90	71	37	105	102	90	76		
	Rang:		12	1	13	9	5	2	6	7	11	14	3	4	7	10		

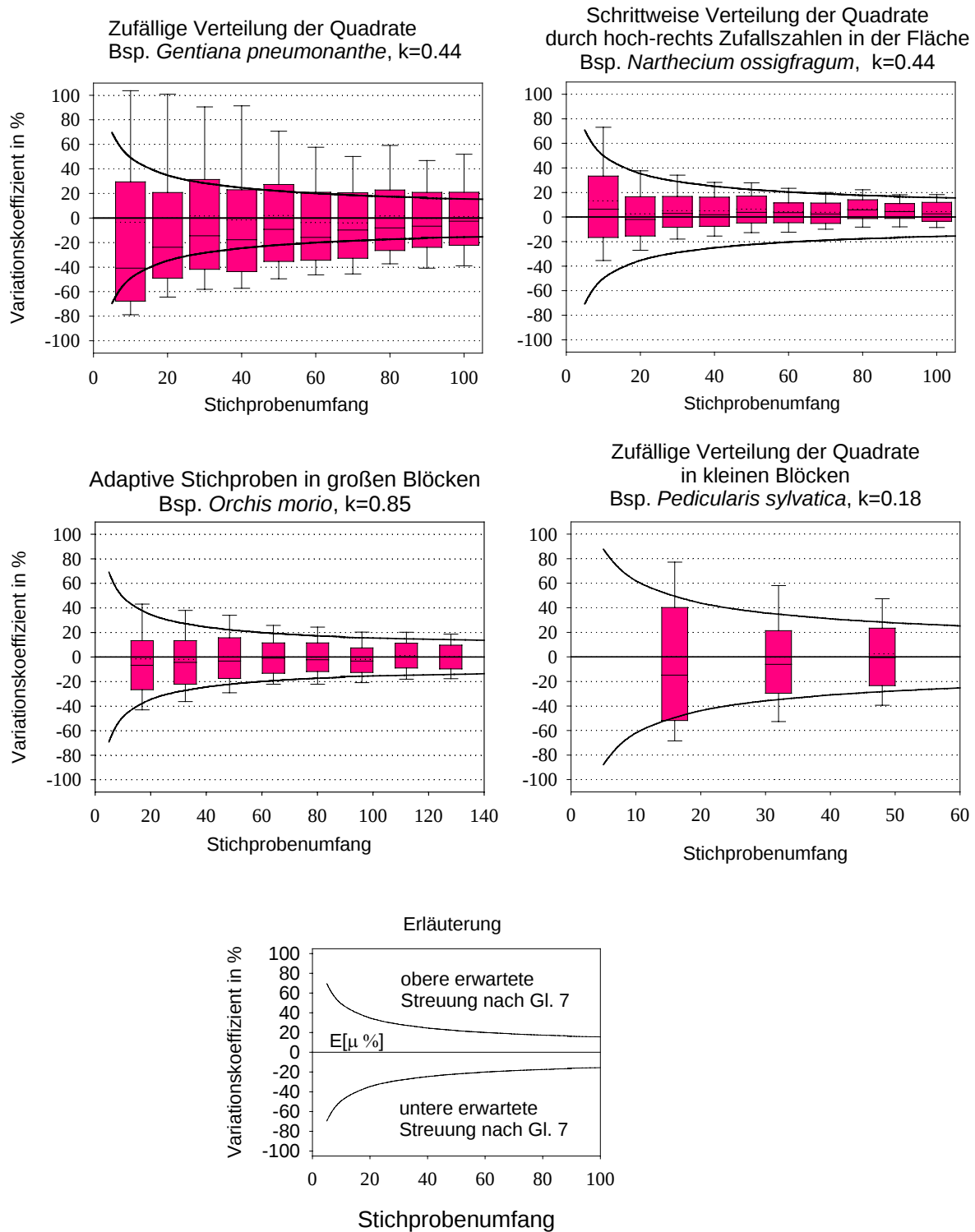


Abb. 5: Beispiele simulierter Dichteschätzungsverfahren hinsichtlich Erwartungstreue und Standardfehler aus 200 Wiederholungen

Mit den gegebenen Kriterien ergaben sich folgende Bonitursummen und Ränge der einzelnen Verfahren (Tab. 5). Das Ergebnis ist ein horizontaler Vergleich zur Beurteilung der Verfahren und ein vertikaler zur Charakterisierung der Standorte/Art.

Es ergibt sich kein einheitliches Bild, das ein bestimmtes Verfahren hervorheben würde. Man kann zwischen sehr großen und mittleren Standardfehlern unterscheiden. Relativ kleine Standardfehler sind die Ausnahme. (Anlage C). Insgesamt unterschreitet in kaum einer Situation der Fehlerbalken, bzw. die 5/95% Quartile des Boxplots den Erwartungswert, d. h. in fast keinem Fall ist eines der Verfahren signifikant besser als der theoretische zu erwartende (Ideal)-Wert. In mehreren Beispielen sind sogar die 50% Quartile größer als die theoretisch zu erwartende Streuung. Nur für die Moorlilie (*Narthecium ossifragum*), an zwei Fundorten mit kleinen k 's vertreten, werden für die Streuung Maximalnoten in der Bonitur vergeben. Aber gerade die Moorlilie ist wegen ihrer Wuchsform (sehr geringer Platzbedarf, hohe Dichten / cm^2) nicht für die Schätzung in 1m^2 -Zählrahmen geeignet. Die kleinräumigen hohen Dichten machen die Auszählung in $1/4$ bzw. $1/8 \text{ m}^2$ -Formaten in der Praxis zwingend erforderlich. Auch die Abgrenzung nach unten zeigt, dass die 50% Quartile der Beobachtungen mehrheitlich über den Grenzwerten liegen, und die schon pessimistischen Erwartungswerte noch überschritten werden. Bei den systematischen Verfahren hat das Hoch-Rechts-Verfahren am besten abgeschnitten.

Auch wenn die zufälligen Stichprobenverfahren bei heterogenen Grundgesamtheiten theoretisch zu den schlechteren Verfahren gehören, zeigt sich in dieser Studie, dass die mit diesem Verfahren erzielten Ergebnisse nicht von den aufwendigeren Verfahren abweichen und deutlich besser sind als alle hier getesteten systematischen Verfahren.

Der kleinste Standardfehler wird durch die geschichteten Anlagen erzielt. Innerhalb dieser Verfahren wurde das zu untersuchende Feld entweder in große Blöcke (durchschnittlich 14-18) oder in kleine Blöcke (durchschnittliche Anzahl 56-80) eingeteilt. Hier führt die Anlage der kleinen Blöcke mit insgesamt vielen Quadraten zu den besten Ergebnissen. Es zeigt sich, dass die 2-Schritt, adaptiven Verfahren zwar mit einem erhöhten Arbeitsaufwand verbunden sind, der Untersucher aber nicht mit einer entsprechenden Gewinn an Genauigkeit belohnt wird. Hinzu kommt, dass bei kleinen Stichproben/Block plus adaptiven Verfahren die Population im Mittel unterschätzt wird. Auch wenn das Verfahren sich auf Rang 2 bzw. 3 der Boniturliste (Tab. 5) befindet, muss bei ungünstigen Verhältnissen mit einem größeren Bias gerechnet werden. Eine Übersicht der Ergebnisse über alle Kombinationen und Verfahren ist in Anhang C dargestellt.

Im vertikalen Vergleich zeigt sich, dass bestimmte Arten an ihrem Standort wesentlich einfacher zu erfassen sind als andere. Auch wenn die Dispersion in allen Beispielen als geklumpfte Verteilung mit ähnlichem k vorliegt, variiert die Größe des Standardfehlers zwischen den Standorten, unabhängig vom Stichprobenverfahren. Als Beispiel für eine schwierige Dichteschätzung sei hier der Lungenenzian (*Gentiana pneumonanthe*) genannt, indem die 50%-Quartile aller Verfahren den theoretischen Erwartungswert überschreiten. Als Beispiel für ein einfaches Problem dient der Nordische Augentrost (*Euphrasia frigida*). Unabhängig vom Verfahren werden kleine, akzeptable Standardfehler erreicht.

Arnica montana wurde in der Analyse ausgelassen, da das spezifische Wuchsverhalten eine Anwendung der Schätzrahmen nicht sinnvoll erscheinen lässt und Alternativverfahren zur Dichteschätzung verwendet werden sollten. Folgende Verfahren wurden zwar vorgestellt, sind aber in der Analyse nicht weiter verfolgt worden.

Variable Schichten: Es gelingt nicht homogene Verteilungen innerhalb der Schichten zu finden, der Standardfehler ist wesentlich höher als in den Blockanlagen.

Das Verfahren mit festen Quadraten, die subjektiv in Bereichen hoher Dichten gelegt werden, hat für den Lungenenzian gute Ergebnisse mit leichter Überschätzung der Population geliefert, für andere Standorte/Arten war das Ergebnis unzufriedenstellend und wurde nicht weiter verfolgt.

Suchpfade sind de facto eine andere Form der systematischen Verfahren, mit allen Vor- und Nachteilen, gelten nur für gut sichtbare Pflanzen. Außerdem wurden in der Praxis diese Verfahren als Totalerhebungen verwendet, und sind daher in der Form nicht mehr als ein Stichprobenverfahren anzusehen.

5 Diskussion

Die hier untersuchten seltenen Pflanzenarten erfüllen i.d.R. die Hypothese, dass ihre räumliche Verteilung als heterogen bis stark geklumpft anzusehen ist. Mit Ausnahme dem Beispiel *Euphrasia frigida* im Perlenbachtal (eine Art mit hoher Dichte und einem k -Wert von 2.4) liegen die gemessenen k -Werte der NGB eher im Bereich von 0.4. Konsequenterweise ist ein sehr großer Stichprobenaufwand für einen vertrauenswürdigen Schätzer notwendig. Die gute Übereinstimmung der simulierten Beobachtungen mit der NGB bedeutet nicht notwendigerweise, dass die natürlichen Po-

pulationen ebenfalls dieser theoretischen Verteilung folgen. Es kann durchaus daran liegen, wie die untersuchten Populationen generiert wurden. Der hier verwendete MT-Prozess kann zu einer NGB führen (BOSWEL & PATIL 1970). Die erfolgreiche Anpassung der NGB verleitet dazu, den Parameter k der NGB als dispersionsspezifischen Parameter zu interpretieren. Das ist nur gültig für ein konstantes Feld/Zählrahmengröße-Verhältnis. Auch wenn ein kleines k auf eine starke Klumpung hinweist, so darf nicht vergessen werden, dass der Parameter dichteabhängig ist und sich mit der Quadratgröße ändert (Richter et al, 1990). Es sei hier am Rande erwähnt, dass k nicht als einziger Dispersionsparameter anzusehen ist, sondern über das Verhältnis Mittelwert zu Varianz andere Dispersionsindices erstellt worden sind (LLOYD 1967; IWA0 1968; PERRY 1981; REED 1983).

Die Auswertung der verschiedenen Stichprobenverfahren widerspricht überraschend den Erwartungen bei Projektbeginn. Mit den hier simulierten Verfahren ist es selbst bei sehr hohen Stichproben nur in Ausnahmen gelungen, einen kleineren Standardfehler zu erhalten, als von der theoretischen Seite zu erwarten gewesen wäre. Zweitens unterscheiden sich die Verfahren anhand der Bonitur mit Ausnahme der systematischen Verfahren nicht oder nur marginal. Der Rang innerhalb der Auswertung spiegelt ein realistisches Bild der Verfahren zueinander wider. Zu unterscheiden sind Teile der systematischen Verfahren, die in keiner Situation eine realistische Bestandsaufnahme liefern (Diagonale, Gräte). Der Vorteil systematischer Verfahren liegt in ihrer einfachen Anlage und Durchführung. Wie hier zu ersehen ist, sind sie im allgemeinen anfällig auf das Muster der räumlichen Verteilung und führen zu verzerrten Schätzern. Gute Ergebnisse in einzelnen Situationen sind eher zufälliger Natur. Die anderen Verfahren, die ebenfalls zur Klasse der systematischen Verfahren gehören, sind im Ergebnis mit einer mittleren Qualität zu beurteilen. Sie zeigen von der Gesamtbonitur her eine schwächere Leistung, aber z.T. sind die 50% Quartile von geringerer Streuung. Das negative Ergebnis ist ebenfalls die Anfälligkeit gegenüber der Dispersion bedingt.

Die zufälligen Verfahren ergeben einen überraschend guten Rang innerhalb dieser Auswertung. Auch wenn diese nicht optimal sind, ist die Anfälligkeit der zufälligen Verfahren auf die Art der Dispersion nicht so ausgeprägt wie bei den systematischen Verfahren.

Am besten schneiden die geschichteten oder Blockverfahren ab. Ein erhöhter Aufwand im Versuchsdesign bewirkt einen kleineren Standardfehler. Zu unterscheiden

ist zwischen den beiden hier getesteten Verfahren. Kleine Blöcke erfüllen eher die Voraussetzung der Blockhomogenität und führt zu kleineren Standardfehlern, während in großen Blöcke diese Voraussetzung in einigen Beispielen nicht erfüllt ist. Daher ist der relative Unterschied von großen Blöcken zu den zufälligen Verfahren in den Graphiken der Anlage C nicht so auffällig und führt zu ähnlichen Boniturwerten. Aber ein numerischer Vergleich der 50%-Quartile der Boxplots zwischen zufälligen und Blockverfahren ergibt durchaus quantitative Unterschiede. So sind am Beispiel des Lungen-Enzians die Boxplots des Blockverfahrens mit großen Blöcke für ein Stichprobenumfang von 80 Quadraten um 14% kleiner als die Boxplots des zufälligen Verfahrens, bei 100 Quadraten beträgt der relative Unterschied gerade mal 3%.

Der relative Vorteil der adaptiven Verfahren ist nicht offensichtlich. Vor allem, wenn einzelne Blöcke mit nur 1 oder 2 Stichproben erfasst werden, werden die genaueren Messungen in den Blöcken mit größeren Stichproben übergewichtet. Im Mittel werden die Population daraufhin unterschätzt.

Interessanter ist der vertikale Vergleich über die Standorte. Unabhängig vom verwendeten Verfahren ergibt sich auch hier eine Abstufung. Anhand der Simulationen kristallisieren sich Standorte heraus, die unabhängig der Wahl eines Verfahrens gute Populationsschätzer ermöglichen, im umgekehrten Fall ist die Varianz so groß, dass innerhalb der gegebenen Stichprobenumfänge kein vertrauenswürdiger Schätzer ermittelt werden kann. Insbesondere die Standorte im zweistelligen Rang stellen den Untersucher vor Probleme, auf die mit verstärktem Aufwand, Planung und arbeitsintensiven Design reagiert werden muss. Man kann 2 Gruppen erkennen. Die Standorte im Rang 1 bis 5 werden durch zufällige Verfahren gut erfasst, während die Standorte mit zweistelligen Rängen zwar allgemein durch einen größeren Standardfehler charakterisiert sind, hier aber die Blockanlage (*stratified sampling*) als die Alternative angesehen werden muss, um einen einigermaßen vertrauenswürdigen Wert für die Populationsdichte zu erhalten.

Sind die gewählten Beispiele repräsentativ für die Gesamtpopulationen des Untersuchungsgebietes, so kann man davon ausgehen, dass in der Hälfte der Untersuchungsstandorte zufällige Verfahren eine ausreichende Genauigkeit liefern würden, während in 30% der Fälle die aufwendigeren Blockverfahren angewendet werden sollten. Sind die betreffenden Standorte auch weiterhin Gegenstand von Untersuchungen, so ist bei den Standorten im zweistelligen Rang auf jeden Fall eine geschichtete Anlage angebracht. Liegen keine Vorinformationen vor, so ist die Blockan-

lage mit kleinen Blöcken das geeignete, aber auch aufwendigste Verfahren. Optimierungen bezüglich des Stichprobenumfangs sind im Bereich des *stratified sampling* möglich. Man muss nicht notwendigerweise aus jedem Block eine Unterstichprobe ziehen, ohne an Genauigkeit zu verlieren. Wie groß die Einsparung sein kann, hängt natürlich wieder von den Populationskenngrößen ab. Daher kann keine allgemein gültige Aussage getroffen werden. Die Ergebnisse verdeutlichen aber, dass die Verwendung von zufälligen Verfahren zwingend erforderlich ist, ansonsten ist der Mittelwert verzerrt und reflektiert nicht die Eigenschaften der Population.

Auch wenn die adaptiven Verfahren nicht zu den gewünschten Verbesserungen geführt haben, sind Vorteile in der Anwendung stärker verschachtelter System möglich (KUNO 1976), die sich noch spezifischer an der räumlichen Heterogenität orientieren und innerhalb der Blöcke größtmögliche Homogenität erzeugen (GREIG-SMITH 1964; RIPLEY 1981) .

Die Technik der Transekte (Suchpfade) ist in diesem Simulationsprojekt nicht weiter verfolgt worden, denn sie gilt i.W. für blühende Pflanzen. In der Praxis wurden sie für Totalerhebungen verwendet.

Die Ergebnisse verdeutlichen, dass eine Bestandsaufnahme mit Hilfe eines Schätzrahmens mit einem vertretbaren Aufwand den hier getesteten Verfahren unter den gegebenen räumlichen Verteilungen nur schwer und mit hohem personellen Aufwand realisierbar sind. Auch das eigentliche Ziel, mögliche Änderungen innerhalb der Populationsdichte statistisch abzusichern, kann als nicht gelöst gelten, da eine Variationsbreite von unter Streuung von 20% in keinem Fall unterschritten wird.

Abschließend sei nochmals ausdrücklich darauf hingewiesen, dass es sich bei den Ergebnissen um Simulationen handelt, in denen versucht wurde, die Realität auf mehreren Ebenen abzubilden. Zu diesen Ebenen gehören einerseits die Populationsbildung in bezug auf Dichte und Klumpung und andererseits die Stichprobenverfahren, die auf idealisierten rechteckigen Feldern mit mehreren gleich- und normalverteilten Zufallszahlen erzeugt wurden. So kann es aufgrund der Zufallszahlenkombinationen durchaus zu völlig unrealistischen Verteilungen der Quadrate kommen, die in der Praxis so nicht durchgeführt würden.

6 Literatur

ANSCOMBE, F.J., 1950: Sampling theory of the negative binomial distribution and logarithmic series distribution, *Biometrika* 37. S. 358 - 382.

BOSWELL, M.T., PATIL, G.P., 1970: Chance mechanism generating the negative binomial distribution. Aus: PATIL, G.P. (Hrsg.), *Random Counts in Scientific Works*. Vol. I, S. 3 - 22.

BROWN, D. & ROTHERY, P., 1993: *Models in Biology: Mathematics, Statistics and Computing*, John Wiley & Sons, 688 S.

BUCKLAND, S.T., ANDERSON, D.R., BURNHAM, K.P. & LAAKE, J.L., 1993: *Distance Sampling, Estimating abundance of biological populations*, Chapman & Hall, 446 S.

COCHRAN, W.G., 1972: *Stichprobenverfahren, Sampling techniques*, de Gruyter, Berlin, New York. Dt. Übersetzung Dr. W. Böing.

DIGGLE, P.J., 1983: *Statistical analysis of spatial point patterns (Mathematics in biology)*. Academic Press, London.

GREIG-SMITH, P., 1952: The use of random and contagious squares in the study of the structure of plant communities. *Annals of Botany* 16. 293 - 316.

IWAO, S., 1968: A new regression method for analysing the aggregation pattern of animal populations. *Researches on population ecology* X. S. 1-20

KUNO, E., 1976: Multistage sampling for population estimation, *Researches on population ecology* 18. S. 39 - 56.

LLOYD, M., 1967: Mean crowding. *Journal of Animal Ecology* 36. S. 1 - 30.

PERRY, J.N., 1981: Taylor's power law for dependence of variance on mean in animal populations. *Applied Statistics* 30. S. 254 - 263.

REED, W.J., 1983: Confidence Estimation of ecological aggregation indices based on counts - A robust procedure. *Biometrics* 39. S. 987 - 998.

RICHTER, O. & SÖNDGERATH, D., 1990, *Parameter Estimation in Ecology, The Link between Data and Models*, VCH, 218 S

RIPLEY, B.D., 1981: *Spatial statistics*. John Wiley & Sons.

SCHEAFER, R.L., MENDENHALL, W., & OTT, L., 1979, *Elementary Survey Sampling*, 2nd Edition, Duxbury Press, 278 S.

SEBER, G.A.F., 1982: *The estimation of animal abundance*. 2. Auflage, Charles Griffin, London.

SEBER, G.A.F., 1986: A review of estimating animal abundance. *Biometrics* 42. S. 267 - 292.

SOUTHWOOD, T.R.E., 1978: *Ecological Methods*. 2. Auflage Methuen: London.

TAYLOR, L.R., WOIWOD, I.P. und PERRY, J.N., 1978: The density dependence of spatial behaviour and the rarity of randomness. *Journal of Animal Ecology* 47. S. 383 - 406.

TAYLOR, L.R., WOIWOD, I.P. und PERRY, J.N., 1979: The negative binomial as a dynamic ecological model for aggregation, and the density dependence of k . *Journal of Animal Ecology* 48. S. 289 - 304,