



Un programme simple de regression non-lineaire pondere adapte aux estimations de biomasse forestiere

Jacques-Eric J.-E. Bergez, J.L. Bisch, Alain A. Cabanettes, Loic L. Pagès

► To cite this version:

Jacques-Eric J.-E. Bergez, J.L. Bisch, Alain A. Cabanettes, Loic L. Pagès. Un programme simple de regression non-lineaire pondere adapte aux estimations de biomasse forestiere. Annales des sciences forestières, 1988, 45 (4), pp.399-412. hal-02719645

HAL Id: hal-02719645

<https://hal.inrae.fr/hal-02719645>

Submitted on 1 Jun 2020

HAL is a multi-disciplinary open access archive for the deposit and dissemination of scientific research documents, whether they are published or not. The documents may come from teaching and research institutions in France or abroad, or from public or private research centers.

L'archive ouverte pluridisciplinaire **HAL**, est destinée au dépôt et à la diffusion de documents scientifiques de niveau recherche, publiés ou non, émanant des établissements d'enseignement et de recherche français ou étrangers, des laboratoires publics ou privés.

Un programme simple de régression non-linéaire pondérée adapté aux estimations de biomasse forestière

J.E. BERGEZ, J.L. BISCH *, A. CABANETTES, L. PAGÈS **

INRA, Station de Sylviculture, Centre de Recherches d'Orléans,
Ardon, F 45160 Olivet

Summary

A simple non-linear weighted regression computer program for forest biomass estimations

The paper presents a computer program destined to fit experimental data through the nonlinear models :

$$Y = a * X1^{\alpha} * X2^{\beta} + b \quad \text{ou} \quad Y = a * X1^{\alpha} + b$$

This kind of model is particularly interesting for forestry biomass estimations (where Y = dry weight ; X1 = diameter ; X2 = height of the tree).

The program takes into account :

- the research of optimal values for both exponents α and β ;
- the possibility to weight the residuals of the regression with a power function of the explanatory data $X_1^{\alpha} * X_2^{\beta}$ or X_1^{α} as the case may be ;
- the calculation of the confidence bounds (level 95 p. 100) for estimated values ;
- the calculation of the error when a such an equation is used to predict an estimation of the total biomass (ΣY) of a population.

The practical interest of the program stands in its flexibility (conversational program with possibilities of choice and a partial presentation of the important intermediate results on the screen — helping to choose the best option) and in the fact that it is a self sufficient program for the PC.

Key words : Computer program, regression, biomass, error, precision, optimization, confidence bound.

Résumé

On présente un programme permettant l'ajustement de données expérimentales aux modèles non-linéaires :

$$Y = a * X1^{\alpha} * X2^{\beta} + b \quad \text{ou} \quad Y = a * X1^{\alpha} + b$$

(*) Adresse actuelle : Office National des Forêts, Division de Mulhouse, 21, rue de l'Est, F 68100 Mulhouse.

(**) Adresse actuelle : INRA, Station d'Agronomie, Domaine Saint-Paul, B.P. 91, F 84140 Montfavet.

particulièrement utile dans le domaine des estimations de biomasse forestière (Y = biomasse totale d'un arbre ; X_1 = diamètre à 1,30 m ; X_2 = hauteur totale). Le calcul intègre à la fois :

- la recherche des valeurs optimales des exposants α et β ;
- la possibilité de pondérer les résidus de la régression par une fonction puissance de la variable explicative $X_1^\alpha * X_2^\beta$ ou X_1^α selon le cas ;
- le calcul des intervalles de confiance (à 95 p. 100) des valeurs estimées ;
- le calcul de l'erreur commise en appliquant ce modèle à l'estimation de ΣY pour une population.

L'intérêt pratique du programme réside dans sa souplesse d'utilisation (abondance du conversationnel, nombreux choix possibles, résultats partiels en ligne pour aides aux décisions) et son autonomie de fonctionnement sur micro-ordinateur compatible PC.

Mots clés : Logiciel de programmation, régression, biomasse, erreur, précision, optimisation, intervalle de confiance.

1. Introduction

Les recherches menées sur l'estimation de la biomasse d'arbres forestiers, relativement récentes à l'INRA (1979), ont créé de nouveaux besoins et renforcé certaines exigences au niveau des modèles statistiques reliant la biomasse totale d'un arbre et l'une ou plusieurs de ses dimensions, modèles dont la finalité est d'estimer la masse sur pied de peuplements entiers.

Traditionnellement, l'estimation de la production forestière était principalement basée, au niveau individu, sur le cubage de gros arbres de futaie, et ne concernait que la dimension volume du tronc arrêté à un diamètre minimum non-nul. Pour des raisons géométriques (prise en compte du seul tronc et forme de celui-ci), ces volumes s'ajustent convenablement à des modèles du type (BOUCHON, 1974) :

$$V = a * D^2 + b \quad (1)$$

$$\text{ou : } V = a * D^2 * H + b \quad (2)$$

par analogie avec la formule de cubage d'un cône :

$$V = f * \pi (D^2/4) * H \quad (3)$$

où V est le volume, D le diamètre à 1,30 m et H la hauteur du tronc, a et b des paramètres, et f le coefficient de forme du tronc, égal à $1/3$ pour un cône parfait. La simplicité des calculs, à une époque où les ordinateurs n'étaient pas sur le marché, a fait passer l'utilisation des modèles (1) et (2) dans la pratique courante.

La prise en compte récente de l'arbre total dans les recherches forestières, et la nécessaire substitution de la grandeur « biomasse » à la grandeur « volume », ont amené à une modification importante de la nature de la grandeur à estimer, dont les propriétés sont plus complexes (présence des branches, variabilité de la densité du bois). Les nouvelles facilités de calcul et, pour certains auteurs, l'intervention d'objectifs plus « explicatifs » (tests d'hypothèses biologiques) (BOUCHON, communication personnelle) ont amené à délaisser les modèles géométriques (1) et (2), même si ces modèles donnent satisfaction dans certains cas (SATO *et al.*, 1982 ; ALEMDAG, 1984 ; LAVIGNE, 1982). Quelques auteurs ont adopté un modèle de type allométrique :

$$B = a * X^\alpha \quad (4)$$

où B est la biomasse, et X peut représenter D^2 ou $D^2 * H$ (WILLIAMS *et al.*, 1984 ; SATOO *et al.*, 1982 ; PASTOR *et al.*, 1984) ; l'ajustement est alors effectué souvent sous la forme linéarisée bi-logarithmique, ce qui régularise la variance résiduelle mais entraîne un biais d'estimation (FLEWELLING & PIENAAR, 1981). L'utilisation directe du modèle allométrique, non-linéaire, grâce à des méthodes itératives, évite ce biais et peut donner de meilleurs résultats que les ajustements linéaires polynômiaux (OUELLET, 1983).

Une adaptation du modèle allométrique, déjà tentée par PAGÈS (1986), a été reprise ici, de manière à s'affranchir de l'hypothèse d'une ordonnée à l'origine nulle :

$$B = a * D^{\alpha} + b \quad (5) \quad \text{ou} : B = a * D^{\alpha} * H^{\beta} + b \quad (6)$$

Ce nouveau modèle non-linéaire correspond à une généralisation des modèles (1) et (2) lorsque la valeur des exposants n'est plus fixée a priori. La non-fixation a priori des valeurs des exposants α et β à 2 et 1 respectivement correspond à un certain nombre de besoins : les diamètres (D) peuvent être mesurés à 1,30 m ou à la base de la tige et l'on constate alors des écarts significatifs entre valeurs de α ; d'autre part, on constate empiriquement que la fluctuation libre de α et β permet d'éviter des ordonnées à l'origine très négatives ; enfin, d'autres applications de ce modèle, notamment pour estimer la biomasse de cépées de taillis (CABANETTES, 1987), justifient cette démarche.

Bien qu'il existe une technique générale d'ajustement de ces modèles non-linéaires (méthode du maximum de vraisemblance avec itérations : BOUVIER *et al.*, 1985), nous avons choisi un mode de calcul plus simple, pouvant fonctionner sur des calculateurs de petite capacité, utilisant la méthode des moindres carrés pondérés, dans lequel les paramètres sont estimés selon 2 phases hiérarchisées, la priorité étant donnée à l'estimation des exposants.

La non-constance de la variance résiduelle par rapport au modèle nous a amenés à intégrer au calcul de régression une pondération des résidus (TOMASSONE *et al.*, 1983) dont la forme choisie, courante pour les biomasses forestières (OUELLET, 1983 ; LAVIGNE, 1984), correspond à une variation de la variance résiduelle selon une fonction puissance de la variable explicative.

Enfin, le calcul de l'erreur d'estimation de la biomasse de la population, très rarement pris en compte par les logiciels statistiques existants, a été associé aux calculs de régression proprement dits.

L'ensemble de ces calculs est réalisé dans le programme REGRE présenté ci-dessous.

2. Description

2.1. Modèle de régression

On peut utiliser les 2 modèles (5) ou (6) selon le nombre de variables explicatives utilisées, le modèle (5) se ramenant à une simplification de (6) en posant initialement $\beta = 0$.

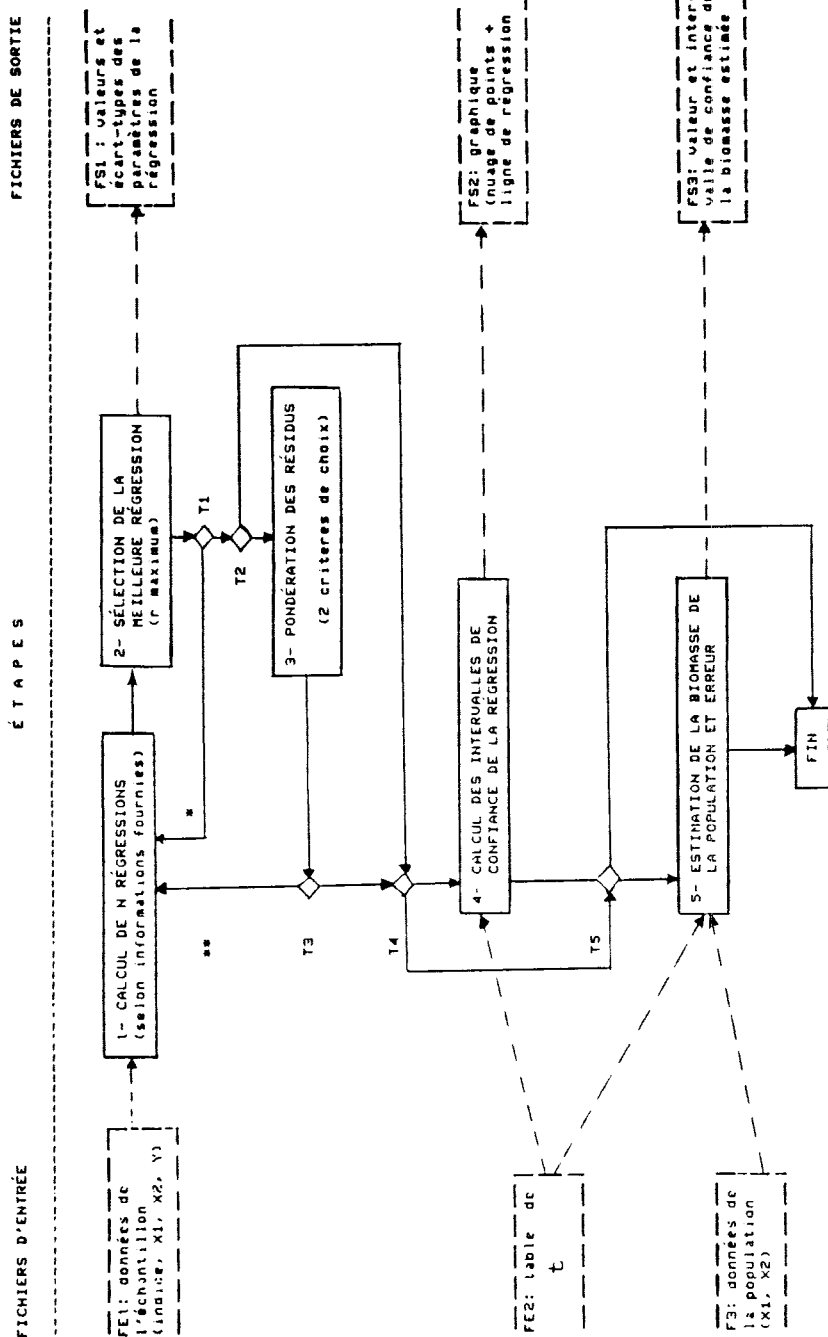


FIG. 1

Organigramme de fonctionnement du programme REGRE.
Running diagram of the REGRE program.

2.2. Principes de fonctionnement

Il y a recherche de la (des) valeur(s) optimale(s) de l'exposant (ou des 2 exposants α et β) avec possibilité de prendre en compte simultanément une pondération des résidus de la régression qui soit fonction de la variable explicative D^a ou $D^a * H^b$ selon que l'on adopte le modèle (5) ou le modèle (6). Cette recherche s'effectue en fixant les exposants à une série (simple ou double) de valeurs et en calculant pour chacune d'elles les paramètres « a » et « b » et le coefficient de corrélation « r » par la méthode des moindres carrés. L'équation fournissant le coefficient de corrélation le plus élevé est retenue. La dépendance existant entre les valeurs optimales des exposants et la pondération choisie nécessite des itérations successives jusqu'à stabilisation des valeurs des paramètres de régression et de pondération.

2.3. Etapes de calcul (fig. 1)

— Etape 1 : le fichier de lecture est celui de l'échantillon ; il doit contenir sur chaque ligne : un indice (permettant de distinguer plusieurs populations), la (ou les deux) variable(s) d'entrée, et la variable à expliquer ; aucun tri préalable n'est nécessaire. La recherche des valeurs optimales de α et β est effectuée entre des bornes et selon un pas d'accroissement fixés par l'utilisateur. Le nombre maximal d'équations de régression calculées est limité (et contrôlé) à 1 000.

— Etape 2 : parmi l'ensemble des équations calculées, seule l'équation fournissant la valeur de « r » maximal est sélectionnée ; les paramètres des autres équations calculées sont néanmoins consultables sur fichier.

— Etape 3 : la pondération des résidus est facultative ; il est toutefois conseillé de ne pas éluder cette étape lors de la première « alternance » entre les étapes 1 et 3, et de la pratiquer ensuite jusqu'à stabilisation des valeurs successives obtenues pour le paramètre « p » décrit ci-dessous. Le coefficient de pondération w de la somme des carrés des écarts est de la forme :

$$w = 1/(X^p) \quad (7)$$

où X est la variable explicative des modèles (5) et (6), et « p » un paramètre à déterminer. Le programme fournit en interactif 2 critères pour estimer « p » : une visualisation du nuage des résidus en fonction de X , et un ajustement du modèle suivant à la variance résiduelle « vares » :

$$\text{vares} = k * X^p \quad (8)$$

(où k est une constante). Afin d'effectuer l'ajustement, des classes d'égale effectif (choisi par l'utilisateur) sont définies, et la variance résiduelle est calculée pour chacune d'elles. L'ajustement est ensuite effectué par la méthode des moindres carrés en passant par la forme linéarisée bi-logarithmique du modèle (8).

— Alternances étape 1 $\leftrightarrow$ étape 3 : l'objectif est la remise en cause des exposants α et β sélectionnés à l'étape 1, compte tenu de la pondération choisie à l'étape 3 (les points expérimentaux n'ayant plus le même « poids » pour le calcul de la régression) et, inversement, la modification de la pondération nécessitée par les nouvelles valeurs de α et β . On arrête les itérations lorsque les valeurs obtenues pour α , β et p sont stabilisées. Le fichier FS1 conserve les résultats de l'ensemble des alternances parcourues.

— Etape 4 : elle concerne la dernière équation de régression sélectionnée (en 1 ou 3), pour laquelle est effectué un calcul d'intervalle de confiance des nouvelles valeurs de Y en des points X non-encore échantillonnés, selon la formule (PERROTTE, 1976) :

$$i.c._X = \pm t * sr * \sqrt{\frac{1}{\sum_{i=1}^n w_i} + \frac{(X - \bar{x})^2}{\sum_{i=1}^n w_i (x_i - \bar{x})^2} + \frac{1}{w_x}} \quad (9)$$

avec X = Valeur de la variable explicative pour l'arbre non-échantillonné.

x_i = Valeur de la variable explicative pour l'arbre n° i de l'échantillon (i entre 1 et n).

$\bar{x}$ = Moyenne pondérée de l'échantillon.

n = Effectif de l'échantillon.

w_i = Facteur de pondération pour l'arbre n° i de l'échantillon.

w_x = Facteur de pondération pour l'arbre non-échantillonné, calculé à partir de la formule (7).

sr = Ecart-type résiduel.

t = Valeur de la variable de Student à n-j degrés de liberté (j pouvant prendre les valeurs 2, 3 ou 4 compte tenu des corrélations entre paramètres).

i.c._x = Demi-largeur de l'intervalle de confiance pour la valeur X.

Une représentation graphique synthétique est ensuite fournie (fichier FS2) figurant simultanément le nuage des points observés (6 populations sont distinguables selon la valeur d'indice), la ligne de régression et les limites de confiance issues de (9) (fig. 3).

— Etape 5 : il s'agit de l'estimation de la biomasse de la population, et de l'erreur commise sur cette estimation en tenant compte de la variabilité de l'échantillon, de l'effectif et de la variabilité de la population, des différences de moyennes entre échantillon et population, et de la pondération adoptée (PERROTTE, 1976). L'intervalle de confiance ainsi calculé est :

$$i.c. \pm t * \sqrt{\text{var}} \quad (10)$$

avec t = valeur du t de Student pour (n - 2) degrés de liberté, et var = variance de l'estimation de la biomasse de la population selon PERROTTE (formule simplifiée dans le cas d'une seule variable explicative) :

$$\text{var} = t_1 + t_2 + t_3$$

0,=,+,x,0,0

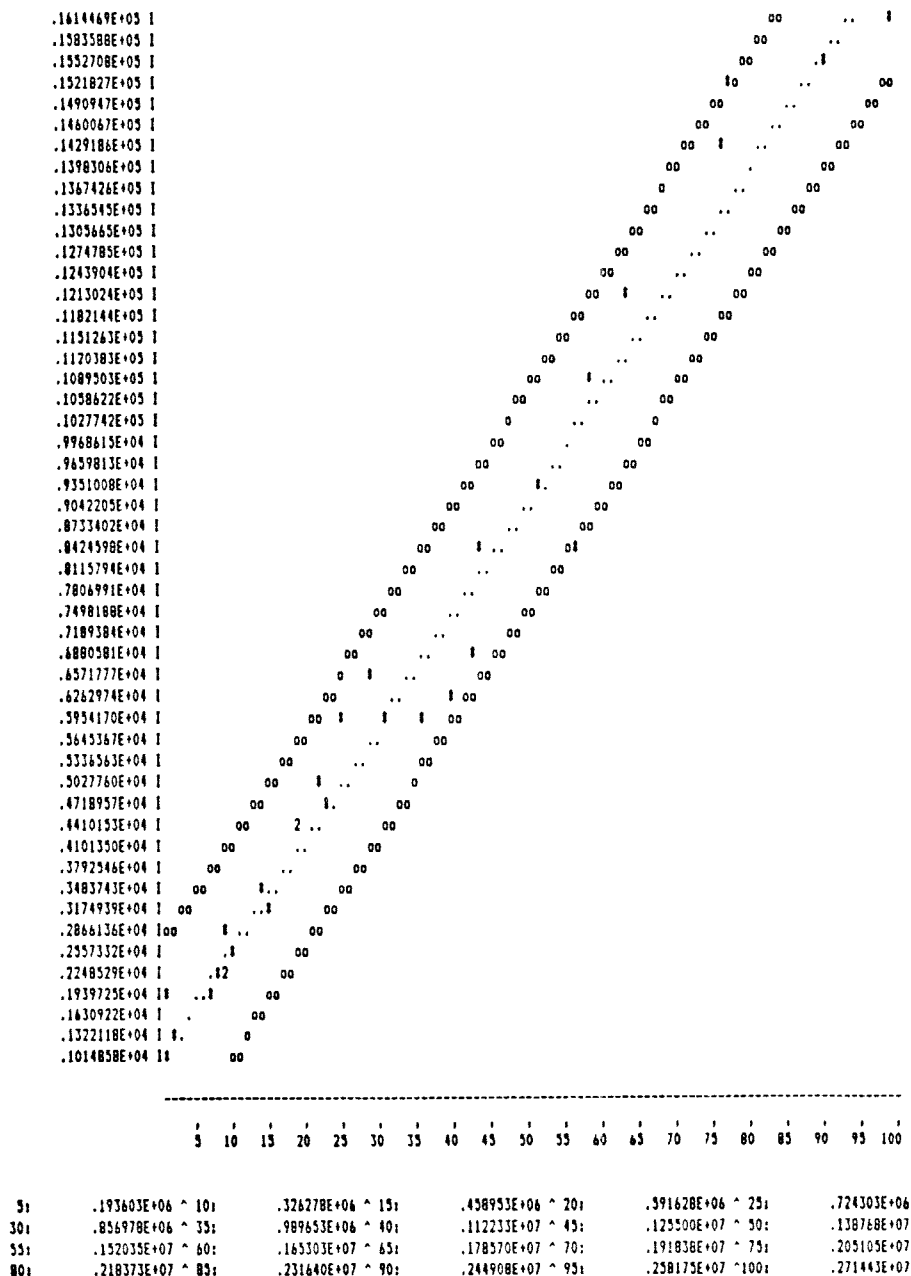


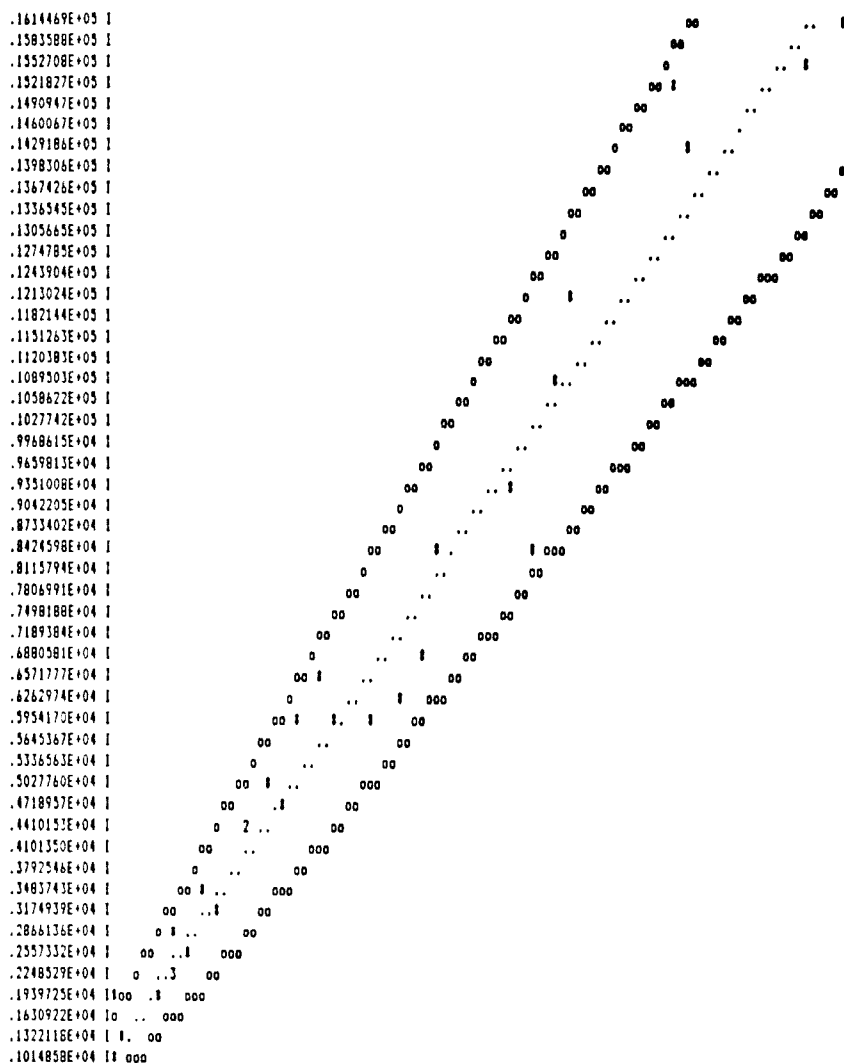
FIG. 3

Exemple de fichiers graphique FS2 représentant les points expérimentaux, la ligne de régression et les intervalles de confiance. A gauche, résultat avant pondération ; à droite, après pondération.

Les données sont issues d'une plantation comparative de provenances d'Epicéa de Sitka (Picea sitchensis Carr.), localisée à Peyrat-le-Château (Haute-Vienne, France) :

BERGEZ et al. (1988). La variable à expliquer (en ordonnées) est la biomasse totale du tronc, les variables d'entrée sont le diamètre à 1,30 mètre et la hauteur totale du tronc.

1,=,+,I,8,4



5 10 15 20 25 30 35 40 45 50 55 60 65 70 75 80 85 90 95 100

5:	.161303E+06 ^ 10:	.276151E+06 ^ 15:	.390998E+06 ^ 20:	.505846E+06 ^ 25:	.620693E+06
30:	.735541E+06 ^ 35:	.850388E+06 ^ 40:	.965236E+06 ^ 45:	.108008E+07 ^ 50:	.119493E+07
55:	.130978E+07 ^ 60:	.142463E+07 ^ 65:	.153947E+07 ^ 70:	.165432E+07 ^ 75:	.176917E+07
80:	.188402E+07 ^ 85:	.199886E+07 ^ 90:	.211371E+07 ^ 95:	.222856E+07 ^ 100:	.234341E+07

Example of graphic file FS2 showing experimental points, regression line and confidence bounds. On the left, without weighting ; on the right, with weighting. Data proceeding from a Sitka spruce (*Picea sitchensis* Carr.) plantation of ten provenances, ten years old, localized in Peyrat-le-Château (Haute-Vienne, France) : BERGEZ et al., 1988. Dependent variable is total woody biomass of the stem (on vertical axis), and independent variables are diameter at breast height and total height of the stem.

avec :

$$t1 = \frac{l^2 * s^2}{\sum_{i=1}^n w_i}$$

$$t2 = \frac{l^2 * s^2 * (\bar{X} - \bar{x})^2}{\sum_{i=1}^n w_i (x_i - \bar{x})^2}$$

$$t3 = s^2 * \sum_{k=1}^l \frac{1}{w_{Xk}}$$

où :

s^2 = Variance résiduelle.

$\bar{X}$ = Moyenne non-pondérée de la population.

$\bar{x}$ = Moyenne pondérée de l'échantillon.

w_i = facteur de pondération de l'arbre n° i de l'échantillon.

w_{Xk} = Facteur de pondération pour l'arbre n° k de la population à estimer, calculé en fonction de X_k à partir de la formule (7).

l = Effectif de la population.

n = Effectif de l'échantillon.

2.4. Caractéristiques techniques

Le programme est écrit en FORTRAN 77 Microsoft (version 3.2.), et est constitué d'un programme principal et de 4 sous-programmes :

- visualisation des résidus ;
- calcul de l'exposant de pondération ;
- calcul des intervalles de confiance et graphique ;
- calcul de la biomasse de la population et de l'erreur sur son estimation ;
- en entrée : 3 fichiers (dont 2 fichiers « utilisateur » et 1 fichier fixe) ;
- en sortie : 3 fichiers « résultats » et 4 fichiers intermédiaires.

Le dimensionnement pour lecture des données est limité à 500 individus maximum, et le nombre maximum de régressions calculées en étape 1 est limité à 1 000 ; ces deux limites fixées par le programme permettent son exécution sur compatible PC 512 Koctets.

3. Utilisation

— Les 2 fichiers FE1 et FE3 sont de format libre (demandé en conversationnel), et doivent contenir, pour le premier, un indice, une ou deux variables d'entrée et une variable à expliquer, et pour le deuxième, la ou les variables d'entrée, les variables

étant lues en format réel ; le fichier FE2, qui contient la table de Student fournissant $t(n)$ au seuil de 5 p. 100 est fixe et fourni avec le programme.

— Un guide conversationnel (questions/réponses) permet d'assister en direct l'utilisateur sur la nature et la forme des informations à fournir, et les possibilités de choix ; en cas d'erreur d'introduction, des retours conditionnels permettent de rectifier des réponses erronées sans sortir du programme (formats de lecture, choix des bornes et du pas de variation des exposants).

— Lorsque l'on n'a aucune idée préalable des valeurs optimales des exposants, il est conseillé de mener leur recherche (phase 1, fig. 1) en 2 temps :

1. localisation grossière à l'aide d'un intervalle de valeurs et d'un pas d'accroissement élevés ;

2. focalisation sur un intervalle de valeurs plus réduit et diminution du pas d'accroissement. Cette stratégie est possible sans sortir du programme grâce au test conditionnel T1 (fig. 1) ; elle permet d'éviter le calcul d'un trop grand nombre de régressions durant la phase 1.

— Lorsque l'on désire au contraire fixer a priori l'une ou les deux valeur(s) des exposants, il suffit de choisir égales les bornes inférieure(s) et supérieure(s) de variation correspondantes.

— Pour l'ajustement de la variance résiduelle (étape 3), il est conseillé à l'utilisateur de choisir un nombre d'individus par classe qui ne soit pas inférieur à 3, tout en s'efforçant d'obtenir un nombre de classes suffisant (au moins supérieur à 5), afin d'avoir des données de variance suffisamment stables et nombreuses, garantissant ainsi a priori un ajustement fiable.

— Enfin, des informations partielles sur les calculs et résultats intermédiaires (paramètres de la meilleure régression, graphe des résidus) sont fournies en cours de programme, à titre de base de décision pour la suite, ou à titre de contrôle (fig. 2).

— Le programme est en libre accès et peut être obtenu sur simple demande. Pour tous renseignements, contacter : Alain CABANETTES, INRA, Station de Sylviculture, Centre de Recherches d'ORLÉANS Ardon, 45160 Olivet, France.

4. Conclusion

Le programme présenté possède un certain nombre d'intérêts pratiques :

— grâce à des calculs effectués sur tableaux indicés, sa rapidité d'exécution demeure satisfaisante même sur de gros fichiers ;

— il n'est pas nécessaire de connaître le fonctionnement d'un autre logiciel associé pour s'en servir, et l'utilisation peut être faite sur tout ordinateur possédant un compilateur FORTRAN 77 Microsoft (version 3.2.) ;

— des itérations prennent en compte la dépendance existant entre les valeurs optimales des exposants et le facteur de pondération ;

— il y a quantification systématique des erreurs d'estimation liées à l'échantillonnage et au modèle d'ajustement et prise en compte de la non-constance de la variance résiduelle, comme le recommande CUNIA (1979) ;

— la préparation des données est très légère (pas de tri, pas de format fixe) ;

— une grande liberté existe au niveau des choix (pondération, variation des exposants) et une aide en ligne est fournie pour les décisions.

Il y a cependant un certain nombre de limites liées au programme :

— le modèle de régression adopté, de type allométrique, est unique ; de même pour le modèle de pondération de type « fonction puissance ». Ainsi, bien que ces modèles, assez généraux, puissent convenir en dehors du seul domaine des estimations de biomasse forestière, il est alors souhaitable d'effectuer au préalable une recherche des meilleures modèles à l'aide de logiciels plus performants (BOUVIER *et al.*, 1985) ;

— lorsque les exposants ne sont pas fixés, les écarts-types de la pente et de l'ordonnée à l'origine (fig. 2) sont sous-estimés compte tenu de la corrélation existant entre les paramètres (dont les estimations ne sont pas simultanées) : les valeurs ne sont alors données qu'à titre indicatif ;

— le nombre maximum d'individus est de 500, et celui des régressions calculées avant sélection de 1 000.

Un certain nombre d'applications pratiques ont déjà été obtenues en ce qui concerne les estimations de biomasse au niveau peuplement (étape 5 ci-dessus) et l'erreur associée : PAGÈS (1986) et BISCH (1987).

Remerciements

Nous remercions vivement Antoine MESSÉAN et Georges PERROTTE pour leurs critiques de fond, ainsi que Max BÉDÉNEAU, qui a réalisé l'adaptation du programme sur micro-ordinateur.

Reçu le 18 décembre 1987.

Accepté le 7 avril 1988.

Références bibliographiques

- ALEMDAG I.S., 1984. Total tree and merchantable stem biomass equation for Ontario hardwoods. *Petawawa Nat. For. Inst. Inf. Rep.* PI-X-46, 1-54 (Canada).
- BERGEZ J.E., AUCLAIR D., ROMAN-AMAT B., 1988. Biomass production of Sitka spruce early thinnings. *Biomass* (soumis pour publication).
- BISCH J.L., 1987. Un exemple de conversion d'une table de production en volume en tables de production en biomasse : le chêne dans le secteur ligérien. *Ann. Sci. for.*, **44** (2), 243-258.
- BOUCHON J., 1974. *Les tarifs de cubage*. Ecole Nationale du Génie Rural, des Eaux et des Forêts, Nancy (France), 57 p. + annexes.
- BOUVIER A., GELIS F., HUET S., MESSEAN A., NEVEU P., 1985. Manuel d'utilisation de CS-NL. INRA, Laboratoire de Biométrie, Jouy-en-Josas (France), 180 p.

- CABANETTES A., 1987. Méthodes d'inventaire dans les jeunes taillis. In : *Compte-rendu de la réunion du Groupe Taillis à Montpellier* (France) les 12 et 13 mars 1987, 7 p.
- CUNIA T., 1979. On tree biomass tables and regression : some statistical comments. In : *Forest resource inventory workshop proceedings*, Frayer W.E. Ed., vol. 2, 14 p., Colorado State University, Fort Collins (U.S.A.).
- FLEWELLING J.W., PIENAAR L.V., 1981. Multiplicative regression with lognormal errors. *For. Sci.*, **27** (2), 281-289.
- LAVIGNE M.B., 1982. Tree mass equations for common species of Newfoundland. *ENFOR Inf. Rep. N-X-213*, 1-40 (Canada).
- OUELLET D., 1983. Equations de prédiction de la biomasse de douze essences commerciales du Québec. *ENFOR Rapp. Inf. LAU-X-62*, 1-29 (Canada).
- PAGÈS L., 1986. Lois de croissance en biomasse du taillis : le robinier dans le Val-de-Loire. *Ann. Sci. For.*, **43** (4), 533-550.
- PASTOR J., ABER J.D., MELILLO J.M., 1984. Biomass prediction using generalized allometric regressions for some Northeast tree species. *For. Ecol. Manage.*, **7**, 265-274.
- PERROTTE G., 1976. *La régression linéaire multiple pondérée en vue de l'application aux calculs des tarifs de cubage et des tables de production*. Doc. Formation Continue, Ecole Nationale du Génie Rural, des Eaux et des Forêts, Nancy (France), 14 p.
- SATOO T., MADGWICK H.A.I., 1982. *Forest Biomass*. Série « Forestry Sciences », Martinus Nijhoff/W. Junk Ed., La Hague, Boston, London, 152 p.
- TOMASSONE R., LESQUOY E., MILLIER C., 1983. *La régression. Nouveaux regards sur une ancienne méthode statistique*. INRA, Actualités scientifiques et agronomiques n° 13, Masson Ed., Paris (France).
- WILLIAMS R.A., MC CLENAHAN J.R., 1984. Biomass prediction equations for seedlings, sprouts, and saplings of ten central hardwood species. *For. Sci.*, **30** (2), 523-527.