

REVUE DE STATISTIQUE APPLIQUÉE

F. BERGERET

Y. CHANDON

Recherche d'une séquence de tests optimale en contrôle de fabrication

Revue de statistique appliquée, tome 43, n° 3 (1995), p. 21-33

http://www.numdam.org/item?id=RSA_1995__43_3_21_0

© Société française de statistique, 1995, tous droits réservés.

L'accès aux archives de la revue « Revue de statistique appliquée » (<http://www.sfds.asso.fr/publicat/rsa.htm>) implique l'accord avec les conditions générales d'utilisation (<http://www.numdam.org/conditions>). Toute utilisation commerciale ou impression systématique est constitutive d'une infraction pénale. Toute copie ou impression de ce fichier doit contenir la présente mention de copyright.

NUMDAM

Article numérisé dans le cadre du programme
Numérisation de documents anciens mathématiques
<http://www.numdam.org/>

RECHERCHE D'UNE SÉQUENCE DE TESTS OPTIMALE EN CONTRÔLE DE FABRICATION

F. Bergeret*, Y. Chandon

Motorola Semiconducteurs S.A.

avenue Eisenhower 31023 Toulouse Cedex

Laboratoire de Statistique et Probabilités URA CNRS 745

Université Paul Sabatier 118 route de Narbonne 31062 Toulouse Cedex

RÉSUMÉ

Dans le cadre du contrôle de fabrication d'un produit, différentes méthodes de réduction des temps de test existent. Deux nouvelles approches sont proposées et comparées. Elles consistent à sélectionner pas à pas un sous-ensemble de tests permettant de détecter les rejets, avec pour critère le coût total de fabrication du produit. Une application sur deux produits est présentée.

Mots-clés : *Arbres de Classification, Choix de Variables, Contrôle de Qualité, Bootstrap.*

SUMMARY

In the fabrication control of a product, several methods have been used to reduce test time. Two new methods are proposed and compared. They select step by step a subset of tests that detects the bad parts and minimizes total manufacturing cost. An application for two products is then presented.

Keywords : *Classification Trees, Choice of variables, Quality Control, Bootstrap.*

1. Introduction

Le contrôle de fabrication sert à s'assurer de la qualité d'un produit avant de le livrer au client. Il peut dans certains cas prendre la forme de multiples tests. Dans cette optique, nous cherchons à éliminer des tests redondants et à trouver le meilleur compromis entre le temps de test et le coût des pièces défectueuses non détectées par le test. Les différentes approches existantes pour diminuer le temps de test sont présentées dans la section 2, nous verrons qu'elles sont efficaces seulement pour une gamme limitée de produits. Pour cette raison nous développerons notre approche qui consiste à sélectionner un sous ensemble de tests qui détectent une forte proportion de rejets. Deux méthodes sont alors proposées, la première, présentée dans

* Ce travail a été réalisé avec une convention CIFRE entre le laboratoire de Statistique et Probabilités de l'Université Paul Sabatier et Motorola Semiconducteurs au Centre Electronique de Toulouse.

la section 3, est adaptée aux produits simples sur lesquels peu de tests sont effectués et dont le rendement (la proportion de pièces bonnes) est élevé. La deuxième méthode, présentée en section 4, s'applique aux produits plus complexes pour lesquels un plus grand nombre de tests sont effectués et dont le rendement est plus faible.

2. Quelques approches

Nous présentons dans cette section différentes méthodes déjà utilisées pour réduire les temps de test.

2.1 Sondage

Il s'agit de tester n pièces choisies aléatoirement parmi les N pièces du lot. A partir de ces n pièces un test est effectué pour vérifier si le rendement du lot est suffisamment élevé. Dans ce cas toutes les pièces sont assemblées. Si l'hypothèse d'un rendement correct est rejetée le lot est alors entièrement testé et les pièces défectueuses ne sont pas assemblées. Une des règles de qualité en vigueur dans l'industrie consiste parfois à déclasser un lot dont le rendement est trop bas, ceci pour éliminer des pièces qui risquent d'avoir un défaut non détectable par le test. Sur la base de cette règle il est possible de déterminer la taille n de l'échantillon pour s'assurer de détecter avec une probabilité donnée un lot à déclasser, ce qui est une méthode classique en contrôle de réception par attribut.

Notons :

- y le rendement du lot,
- y_0 le rendement en dessous duquel le lot est entièrement testé. En pratique ce rendement peut être choisi comme étant la limite au dessous de laquelle il est plus économique de tester le lot entièrement que de laisser passer les rejets à l'assemblage,
- y_1 le rendement seuil de déclassement du lot,
- x le nombre de pièces bonnes pour l'échantillon,
- l le seuil de déclenchement du test complet, à déterminer.

Le test est le suivant :

- $H_0 : y \geq y_0$,
- $H_A : y < y_0$.

La règle de décision est la suivante : on accepte H_0 si $x \geq l$.

Pour un risque de 1^{ère} espèce α fixé et en se donnant une certaine probabilité $1 - \beta$ de détecter un lot à déclasser on obtient les deux relations suivantes :

$$P(x < l / y = y_0) = \alpha$$

$$P(x < l / y = y_1) \geq 1 - \beta$$

En utilisant une approximation Normale de la loi Binomiale, justifiée en pratique si la taille de l'échantillon dépasse 100, on en déduit deux relations entre n et l qui donnent la taille de l'échantillon et le seuil de déclenchement du test complet.

Nous avons étudié l'impact d'un test par sondage sur un produit à rendement élevé : le risque de 1^{ère} espèce est fixé à 0,5 % et la probabilité de détecter un lot à déclasser est fixée à 99,5 %. Connaissant la taille N du lot et après les calculs qui donnent la taille n de l'échantillon, on détermine le taux de sondage n/N . Il est de 1/38^{ème} dans notre exemple avec une augmentation du taux de rejets non détectés, estimée sur plusieurs lots de production, inférieure à 2 %. Le bilan financier est globalement très positif mais l'utilisation de cette méthode n'est envisageable que pour une gamme de produits à rendement très élevé, ceci afin de limiter les rejets non détectés. D'autres solutions devront donc être proposées pour pouvoir analyser une gamme de produits la plus large possible.

2.2 Prévision du résultat du test

Une autre approche consiste à tester un certain nombre de pièces et à utiliser les résultats pour prédire la valeur du test sur les pièces suivantes. Kim et Nanthakumar (1991) proposent la méthode suivante :

m pièces sont testées et on note p_b la proportion de pièces bonnes. Si cette proportion dépasse un seuil donné p_0 la pièce suivante n'est pas testée et le résultat du test pour cette pièce est prévu par un modèle de régression estimé à partir des m pièces précédentes. A l'inverse si la proportion p_b est inférieure au seuil p_0 la pièce suivante est testée. Cette procédure est répétée sur k pièces supplémentaires jusqu'au point où la décision d'accepter ou de rejeter le lot entier est prise. Le lot est accepté si la proportion de pièces bonnes calculée sur $m + k$ pièces est supérieure à p_0 et entièrement rejeté sinon. Le choix de m résulte d'un calcul et minimise le temps total de test. Le choix de k résulte d'un calcul semblable à celui effectué pour le sondage et minimise le risque d'erreur.

Cette méthode, comme le sondage, n'est valable que pour des produits à haut rendement et pour des tests relativement longs. En effet elle n'est efficace que si le temps calcul utilisé pour la prévision et la mise à jour est inférieur au temps de test réel.

2.3 Test adaptatif

Marshall (1987) propose une stratégie de test adaptée à la qualité du lot. On distingue deux groupes de tests : les tests dits fonctionnels, qui captent une majorité de rejets, sont systématiquement effectués. Les tests dits paramétriques, pour lesquels le taux de rejet est faible, sont effectués sur une fraction des pièces. La méthode est adaptative : on commence à tester les pièces entièrement, si i pièces consécutives sont bonnes les tests paramétriques ne sont effectués que sur une pièce sur deux (taux $f = 1/2$). À nouveau si i pièces consécutives sont bonnes le taux passe à $1/4$ puis à $1/8$ qui est le taux limite. Si pendant le processus des rejets paramétriques sont trouvés on teste k pièces à 100 %, si elles sont toutes bonnes le taux reste le même sinon on passe à un taux plus élevé. L'application de cette méthode sur une certaine

gamme de produits permet d'obtenir un gain de temps de l'ordre de 30 % avec une variation non significative du nombre de rejets non détectés.

3. Sélection hiérarchique des tests

Dans cette section nous proposons une méthode pour sélectionner un sous-ensemble de tests qui, effectués sur toutes les pièces, permet de capter une majorité de rejets. Cette méthode est basée sur le principe des arbres de classification introduits par Breiman *et al.* (1984). Nous présentons d'abord le cadre dans lequel nous avons construit cette méthode, puis la méthode elle-même.

3.1 Un problème concret d'optimisation de test

Les circuits intégrés sont fabriqués sur des groupes de plaquettes de silicium, chaque plaquette étant constituée de plusieurs centaines de puces. Après un cycle de fabrication de plusieurs semaines, une première série de tests est effectuée, le test sous pointe ou probe. Un ensemble de mesures électriques, dont le nombre peut atteindre mille, est réalisé sur chaque puce de chaque plaquette. Toute puce qui n'est pas conforme aux spécifications est identifiée puis rejetée. Les puces ayant passé le probe sont assemblées dans un boîtier puis testées une seconde fois lors du test final. Le circuit, s'il passe ce test, est alors livré au client.

Le probe sert essentiellement à détecter les mauvaises puces avant l'assemblage car le coût d'un rejet détecté au test final, incluant le prix du boîtier, est plus élevé. Dans toute la suite de l'étude nous aurons donc une vision économique du probe et nous chercherons à optimiser son coût compte tenu des économies qu'il permet de réaliser au test final. Pour cela nous définissons tout d'abord un coût du probe proportionnel au temps de test par plaquette. Cette hypothèse est réaliste si la capacité libérée sert à tester «gratuitement» d'autres produits, ce qui est le cas quand la production est en augmentation constante. Nous verrons dans la suite que toute diminution du nombre de tests sous pointe peut entraîner une augmentation des rejets détectés au test final. Le coût induit se compose de deux parties, le coût d'assemblage et le coût du test final. Nous définissons finalement l'économie ou gain par plaquette comme étant la différence entre l'économie sur le coût du probe et le coût des pertes additionnelles au test final. Maximiser l'économie par plaquette est le but recherché, ce qui équivaut à minimiser le coût total de fabrication et constitue donc un enjeu financier important car la part du test dans le coût total devient de plus en plus élevée.

$\{W_1, \dots, W_n\}$ constitue un échantillon statistique de plaquettes. Chaque plaquette W_i est constituée d'un ensemble de puces qui ne peuvent être supposées indépendantes en raison de fortes liaisons dues à la fabrication alors que les plaquettes sont tirées indépendamment dans différents lots de fabrication. Nous adoptons alors la méthode suivante :

- construction d'une séquence réduite S_i optimale pour la plaquette W_i ,
- évaluation de chaque séquence S_i sur les $(n - 1)$ autres plaquettes, ce qui donne n gains moyens $G_i, 1 \leq i \leq n$,
- choix de la séquence S^* qui maximise le gain moyen.

Cette méthode a deux avantages : elle permet d'avoir un estimateur non biaisé du gain moyen induit par chaque séquence de test réduite puisque calculé, comme en validation croisée, sur les plaquettes qui n'ont pas servi à la construction de la séquence. D'autre part le coût du test est défini de façon naturelle par plaquette.

3.2 Arbres de Classification

L'origine de ces méthodes remonte aux techniques de segmentation dont on trouvera une approche synthétique dans Baccini et Pousse (1975). Breiman *et al.* (1984) ont jeté les bases des arbres de décision binaire dont on trouvera une application détaillée dans Gueguen et Nakache (1988). La méthode nécessite des calculs intensifs sur ordinateur et de nombreuses adaptations ont été proposées, par exemple dans Loh et Vanichsetakul (1983) qui utilisent de façon récursive l'analyse discriminante linéaire pour construire très rapidement un arbre de régression. Une généralisation est proposée par Ciampi (1991) avec des hypothèses sur la distribution de la variable à expliquer et des tests pour construire et élaguer l'arbre. Un exemple d'arbre de régression basé sur le modèle linéaire généralisé est présenté dans Ciampi *et al.* (1990). Cette approche ne sera toutefois pas développée dans la suite car plus éloignée de notre problème. Du point de vue du vocabulaire on parle d'arbre de régression quand la variable à expliquer est quantitative et d'arbre de classification quand elle est discrète. Nous étudierons ce deuxième cas. Le but d'un arbre de classification est double : Il s'agit d'une part de sélectionner les explicatives les plus informatives et d'autre part de prédire la variable à expliquer. Les explicatives peuvent être soit continues soit discrètes soit les deux. L'échantillon initial est divisé successivement en deux sous-ensembles appelés nœuds, chaque nœud étant défini par la réponse à une question concernant une des variables explicatrices. Les divisions s'achèvent quand on atteint un nœud dit terminal, à ce niveau on assigne ce nœud à une classe. A chaque subdivision, la question retenue est celle qui minimise un critère appelé impureté de l'arbre. On peut donc voir un arbre de classification comme un algorithme ayant pour but d'optimiser un certain critère et fonctionnant de façon itérative, c'est l'approche que nous utiliserons dans la suite de l'exposé. Notons enfin qu'après la construction d'un grand arbre, celui ci est élagué afin obtenir des estimateurs non-biaisés du taux d'erreur.

3.3 Les critères

Toute procédure de discrimination est basée sur l'optimisation d'un critère. Dans cette section nous présentons les deux critères utilisés pour sélectionner les tests.

3.3.1 Critère simple

Le critère retenu ici est le nombre de puces non détectées au probe. Au niveau de la plaquette W_i on adopte la méthode suivante :

- sélection du test (de la mesure) qui capte le plus de rejets,
- pour les puces qui ont passé ce test sélection du test qui capte le plus de rejets,
- ainsi de suite jusqu'à détection de tous les rejets de la plaquette.

Il en résulte une séquence de test réduite par rapport à la séquence de test initiale. Notons que pour des raisons liées à l'électronique il est parfois nécessaire d'effectuer un ou plusieurs tests *avant* de pouvoir effectuer un ou plusieurs autres tests. L'algorithme tient compte de ces contraintes de la façon suivante : si le test t_1 doit être effectué avant le test t_2 alors l'introduction du test t_2 est équivalente à l'introduction du test t_2 et du test t_1 .

Ce problème correspond à un cas particulier d'arbre de classification dans lequel la variable à expliquer est binaire (la puce est bonne ou non) et les variables explicatives sont également binaires (le test est bon ou non). Le problème de la reconnaissance d'un chiffre traité dans Breiman *et al* (1984) est proche de ce cas mais avec une variable à expliquer à 10 modalités. L'application de cette méthode permet de réduire de façon significative le nombre de tests mais ne prend pas en compte les temps de test dans la construction de la séquence réduite, pour cette raison nous utilisons dans la section suivante un critère de coût.

3.3.2 Critère de coût

Le coût C_o d'un rejet non détecté au probe est connu, il se compose de deux parties :

- la valeur ajoutée au produit par l'assemblage, C_a
- le coût du test final, C_{tf}

Par définition $C_o = C_a + C_{tf}$. Nous faisons ici l'hypothèse que tout rejet non détecté au probe sera détecté au test final. Définissons :

- une séquence de test t est un vecteur de dimension égale au nombre total q de tests, de coordonnées $t_i = 0$ si le test i est absent et $t_i = 1$ sinon,
- $ND(t)$ est le nombre de rejets par plaquette non détectés par la séquence de test t ,
- $CP(t)$ est le coût du probe par plaquette pour la séquence t ,
- T est la séquence de test complète : $T_i = 1, 1 \leq i \leq q$,
- le gain par plaquette avec la séquence t :

$$G(t) = CP(T) - CP(t) - C_o ND(t).$$

L'objectif est de trouver la séquence t^* qui maximise le gain par plaquette. L'algorithme opère séquentiellement en sélectionnant à chaque étape le test le plus économique dans la détection des rejets et s'arrête quand l'addition d'un test n'améliore plus l'économie réalisée.

3.4 La contrainte de rendement au test final

L'algorithme précédent recherche la combinaison de tests qui offre le meilleur compromis entre coût du probe et coût des rejets non-détectés. Ainsi, pour des produits

dont le coût du probe est très élevé, la séquence retenue risque d'être très réduite avec une baisse de rendement au test final non négligeable. En outre, pour des raisons de stratégie industrielle, le test final s'effectue souvent dans une usine différente de celle où a été effectué le test sous pointe. Le coût des pertes au test final est donc supporté par une usine qui n'a aucune contrepartie, si ce n'est la possibilité d'être livrée plus vite. D'autre part, en cas de dérive du rendement au test final, l'information n'arrive que plusieurs semaines après, ce qui risque de provoquer des problèmes de livraison. Pour minimiser le coût total en limitant les pertes au test final une adaptation de l'algorithme est donc proposée. Il s'agit d'augmenter artificiellement le coût du rejet non détecté. Formellement on cherche la séquence t^* qui vérifie :

$$t^* = \text{Argmax}\{CP(T) - CP(t) - C_o(ND(t) + \lambda)\}$$

λ étant une pénalité positive à fixer par l'utilisateur. Cette pénalité est introduite par analogie avec un programme d'optimisation sous contrainte pour lequel le temps de test noté t serait une variable continue :

$$\begin{cases} \text{Max } \{CP(T) - CP(t) - ND(t)C_o\} \\ \text{sous la contrainte } ND(t) \leq \text{seuil} \end{cases}$$

Le Lagrangien de ce problème s'écrit :

$$L = CP(T) - CP(t) - ND(t)C_o - \lambda(ND(t) - \text{seuil})$$

et ne diffère de notre expression que par une constante indépendante de t .

En pratique λ est progressivement augmenté de 0 à une valeur maximale $\lambda_{\max}$ fixée de façon à ce que le coût corrigé d'un oubli dépasse le coût total de test de la plaquette, $CP(T)$. Pour cette valeur de λ le test sélectionné à chaque étape sera celui qui détecte le plus de rejets, indépendamment du temps de ce test. On retombe sur le critère simple et toute augmentation de λ au delà de $\lambda_{\max}$ ne changera pas la séquence optimale obtenue.

3.5 Intervalle de confiance de l'économie moyenne

Avoir un intervalle de confiance de l'économie réalisée plutôt qu'une estimation ponctuelle est d'un grand intérêt pour pouvoir mettre en œuvre les résultats proposés. Sans hypothèse sur la distribution des gains, il faut utiliser des techniques de rééchantillonnage dont on trouvera une revue dans Efron (1982). Notre estimation du gain est $G^* = \text{Max}(G_1, \dots, G_n)$ et l'utilisation du bootstrap est justifiée car G^* est invariante par rapport à l'ordre de $(W_1, \dots, W_n)$. L'algorithme d'estimation de la distribution bootstrap est décrit ci-dessous :

- tirage avec remise d'un échantillon bootstrap de n plaquettes parmi les n plaquettes de l'échantillon initial,
- recalcul des n gains moyens pour les n séquences optimales,
- choix du gain maximal G_b^* .

Ce calcul est répété un nombre B de fois. En pratique B est le plus grand possible compte tenu du temps calcul disponible. L'intervalle de confiance à un niveau α est obtenu de façon classique : une proportion $\alpha/2$ des gains G_b^* sont inférieurs à la borne inférieure G_l . La borne supérieure G_u est calculée de façon symétrique.

Le bootstrap permet également d'estimer le biais de G^* :

$$\widehat{\text{Biais}} = \left(\sum_{b=1}^B G_b^* \right) / B - G^*$$

Une application du bootstrap est proposée dans la section suivante.

3.6 Application

Le produit analysé par cette méthode a les caractéristiques suivantes :

- 19 tests, 1.6 seconde de test par puce,
- 1790 puces par plaquette, soit 55 mn de test en tenant compte du temps de passage d'une puce à l'autre et du temps d'attente en début de plaquette,
- les temps de test varient dans des proportions de 1 à 10.

L'échantillon contient 30 plaquettes collectées au hasard de la production pendant deux mois. Un programme écrit en Pascal calcule la séquence optimale :

- 1 test sur 19 est retenu,
- la plaquette serait testée 4 fois plus vite,
- la perte de rendement moyenne au test final est estimée à 2 %,
- l'économie moyenne est estimée à 8 % du coût total de fabrication.

L'intervalle de confiance de l'économie réalisée, en pourcentage du coût total, est (7,85% – 8,15%) pour un coefficient de sécurité de 95 % et l'estimation du biais de notre estimateur est légèrement positive. Notons qu'on ne recalcule pas une combinaison de tests à chaque étape et les calculs effectués pour $B = 10000$ sont donc très rapides.

La solution proposée implique une baisse de rendement au test final jugée trop élevée. L'adaptation de la section 3.4 permet de choisir une combinaison entraînant une baisse de rendement au test final moins importante, estimée à 0,5 % (figure 1). Chaque point de la courbe correspond à une combinaison optimale de tests pour une valeur donnée de la pénalité λ . Le point le plus à droite de la courbe correspond à $\lambda = 0$ pour lequel l'économie par plaquette est maximale. Le point le plus à gauche est celui pour lequel le programme cherche à détecter tous les rejets indépendamment des coûts. Entre les deux l'utilisateur choisit un point correspondant à une baisse de rendement au test final supportable pour un gain légèrement moindre qu'à l'optimum.

Une séquence réduite à trois tests entraînant une baisse de rendement au test final estimée à 0,5 % est actuellement en production. Elle permet de tester une plaquette trois fois plus vite.

4. Sélection Pas à Pas

La méthode présentée dans la section précédente est efficace pour des produits relativement simples, c'est-à-dire pour lesquels le nombre de tests par puce n'est pas trop important. Pour des circuits plus complexes la séquence optimale est trop variable d'une plaquette à l'autre et la meilleure séquence sur l'ensemble de l'échantillon n'est pas satisfaisante. Deux changements sont donc proposés pour pouvoir analyser ces circuits :

- la séquence optimale est directement calculée sur l'ensemble des plaquettes de façon à prendre en compte tout l'échantillon dans le calcul de la combinaison de tests la plus économique,
- une procédure permet à chaque étape d'éliminer des tests.

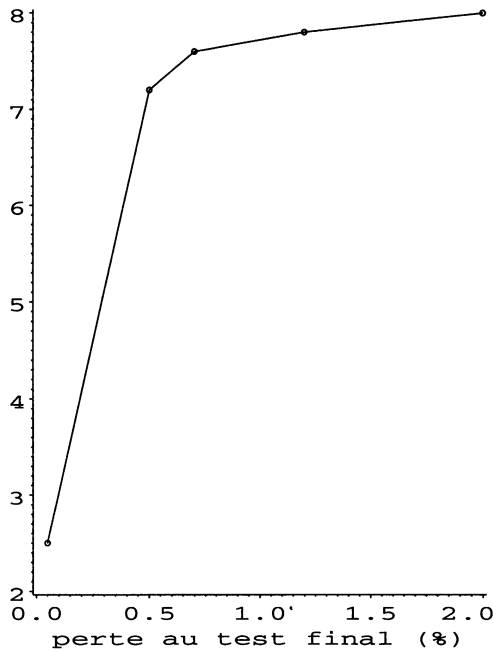


FIGURE 1

*Economie nette (en % du coût total de fabrication)
en fonction de la baisse de rendement au test final.*

4.1 Méthode pas à pas

Le critère reste le même que celui présenté en 3.3.2 mais une amélioration de l'algorithme a été mise en place. La séquence est construite comme précédemment avec, à chaque étape, la possibilité d'enlever un des tests déjà sélectionnés. Cela

permet d'éliminer des tests redondants qui ne l'étaient pas lors de leur introduction. L'algorithme est donné ci-dessous :

- soit t une séquence de test,
- soit S_t l'ensemble des séquences voisines de t . Une séquence t' est dite voisine de t si elle est obtenue à partir de t en ajoutant ou en supprimant un test (ou plusieurs s'il y a des contraintes sur le test en question),
- soit t_0 la séquence de test initiale. En pratique t_0 est la séquence qui ne comporte aucun test.

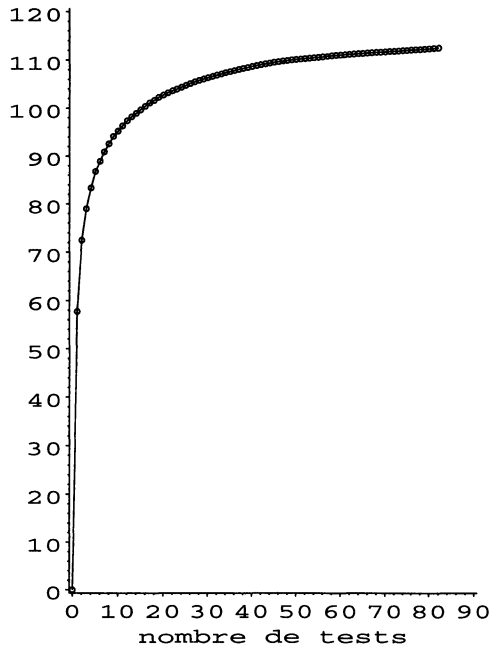


FIGURE 2

Nombres de puces défectueuses détectées en fonction du nombre de tests.

Procédure pas-à-pas;

Début

$t := t_0$;

Répète

Génère($t' \in S_t$);

Si $\text{Eco}(t') \geq \text{Eco}(t)$ alors $t := t'$;

Jusqu'à ($\text{Eco}(t') \leq \text{Eco}(t)$ pour tout $t' \in S_t$);

Fin.

4.2 Estimation du biais

La séquence optimale étant obtenue globalement à partir des n plaquettes de l'échantillon, l'estimation de la baisse de rendement au test final induite va être trop optimiste. La figure 2, construite à partir d'un des produits étudiés, montre en fonction du nombre de tests sélectionnés le nombre moyen de mauvaises puces détectées sur l'échantillon. Il est raisonnable de penser que les premiers tests doivent être retenus quelque soit l'échantillon. En revanche le choix des autres tests permettant de limiter la baisse de rendement au test final sur l'échantillon risque d'introduire un biais et une variance importante dans l'estimation de l'économie moyenne. Ceux et Turlot (1989) passent en revue les différentes méthodes pour estimer le taux d'erreur d'une règle de décision en discrimination : échantillon-test, bootstrap, jackknife, validation croisée. Les techniques de rééchantillonnage permettent une utilisation maximale des données mais sont très coûteuses en temps calcul. Dans notre cas il faudrait recalculer une séquence optimale un grand nombre de fois, ce qui est en pratique beaucoup trop long. Pour cette raison l'estimation du taux de rejet au test final se fait à l'aide d'un échantillon-test ce qui nécessite la collecte d'un nombre plus important de plaquettes. Une proportion des plaquettes, l'échantillon d'apprentissage, sert à construire la combinaison de test optimale et les plaquettes restantes qui constituent l'échantillon-test servent à estimer le taux de rejet. Soit :

- $\hat{r}$ l'estimation de la baisse de rendement test final sur l'échantillon d'apprentissage,
- r^* l'estimation sur l'échantillon-test,

l'estimation du biais de $\hat{r}$ est :

$$\widehat{\text{Biais}} = r^* - \hat{r}.$$

On en déduit immédiatement une estimation du biais de l'estimateur de l'économie moyenne par plaquette. Notons qu'une optimisation de l'algorithme de recherche de la combinaison optimale et l'utilisation d'estimations bootstrap plus efficaces préconisées par Efron (1990) sont à l'étude et devraient permettre d'estimer le biais sans échantillon-test. Il ne sera toutefois pas possible d'obtenir un intervalle de confiance.

4.3 Application

Le produit analysé dans cette section est plus complexe que celui présenté dans la section 3.6, en particulier son rendement est plus faible. Ses caractéristiques sont :

- 94 tests, 2.1 seconde de test par puce,
- 1107 puces par plaquette, soit 44 minutes de test par le même calcul que pour le produit précédent,
- les temps de test varient dans des proportions de 1 à 15.

L'échantillon contient 42 plaquettes collectées au hasard de la production pendant deux mois. La séquence optimale en terme de coût comporte 16 tests, elle est

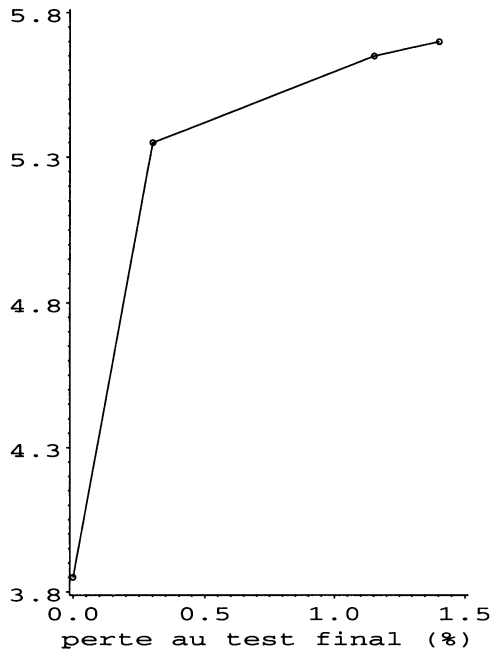


FIGURE 3

Economie nette (en % du coût total de fabrication) en fonction de la baisse de rendement au test final, sur l'échantillon d'apprentissage.

construite sur 30 plaquettes mais entraîne une baisse de rendement au test final estimée à 1.4%. Cette perte étant jugée trop importante, une autre solution est retenue à partir de la figure 3 :

- 24 tests sur 94,
- la plaquette est testée 2,5 fois plus vite,
- la baisse de rendement test final est estimée à 0,3% sur l'échantillon d'apprentissage et à 0,5% sur l'échantillon-test de 12 plaquettes,
- l'économie moyenne, corrigée du biais, est estimée à 5,3% du coût total de fabrication.

5. Conclusion

Pour les produits sur lesquels elle a été appliquée, la méthode pas à pas fonctionne bien compte tenu des économies potentielles qu'elle permet de réaliser. Elle est également applicable pour des puces plus simples mais les résultats sont sensiblement identiques à ceux obtenus par la sélection hiérarchique, vu le faible

nombre de test à effectuer. En revanche, pour des puces très complexes avec un très grand nombre de tests, l'algorithme présenté en 4.1 risque de converger vers un optimum local. Un algorithme d'optimisation de la séquence de test par recuit simulé est à l'étude, il devrait permettre d'approcher l'optimum global pour des produits ayant jusqu'à 10^{150} combinaisons de tests possibles.

Références

- BACCINI A. et POUSSE A. (1975), Segmentation aux moindres carrés : un aspect synthétique, *Revue de Statistique Appliquée*, **XXIII** (3), 17-35.
- BREIMAN L., FRIEDMAN J.H., OLSHEN R.A. and STONE C.J. (1984), *Classification And Regression Trees*, Waldsworth, Belmont CA.
- CELEUX G. et TURLOT J.C. (1989), Estimation de la qualité d'une règle discriminante, *La revue de MODULAD*, 4, 37-44.
- CIAMPI A. (1991), Generalized Regression Trees, *Computational Statistics and Data Analysis*, **12**, 57-78.
- CIAMPI A., LIN Q. and YOUSIF G. (1990), GLIMTREE : RECPAM tree with the Generalized Linear Model, in *Compstat'90, Proceeding in Comp.Stat.*, Physica-Verlag, 21-26.
- EFRON B. (1982), *The Jackknife, the Bootstrap and Other Resampling Plans*, Society for Industrial and Applied Mathematics.
- EFRON B. (1990), More efficient Bootstrap calculations, *Journal of the American Statistical Association*, **85**, 79-89.
- GUEGUEN A. et NAKACHE J.P. (1988), Méthode de discrimination basée sur la construction d'un arbre de décision binaire, *Revue de Statistique Appliquée*, **XXXVI** (1), 19-38.
- KIM H. et NANTHAKUMAR A. (1991), Dynamic Predictive Testing Based on Regression Analysis in Manufacturing Process, *First International Workshop on The Economics of Design and Test*, MCC, Austin, Texas.
- LOH W. and VANICHSETAKUL N. (1983), Tree Structured Classification via Generalized Discriminant Analysis, *Journal of the American Statistical Association*, **83**, 715-728.
- MARSHALL N. (1987), Motorola optimise ses méthodes de test sous pointe, *Electronique Industrielle*, **129/01**, 69-72.