

La sélection grâce à l'utilisation d'une mesure prédictive : une exploration des méthodes

Louis Laurencelle

Volume 28, Number 1, 2005

URI: <https://id.erudit.org/iderudit/1087726ar>

DOI: <https://doi.org/10.7202/1087726ar>

[See table of contents](#)

Publisher(s)

ADMEE-Canada - Université Laval

ISSN

0823-3993 (print)

2368-2000 (digital)

[Explore this journal](#)

Cite this article

Laurencelle, L. (2005). La sélection grâce à l'utilisation d'une mesure prédictive : une exploration des méthodes. *Mesure et évaluation en éducation*, 28(1), 19–36. <https://doi.org/10.7202/1087726ar>

Article abstract

There are but few textbooks that deal with selection and selection procedures, particularly with the use of a predictive measure. Be it for the recruiting or reassignment of personnel, admission into educational programmes, qualification in sports, even quality control for factory-made products, many procedures are available whose mathematical bases are little known or publicized. In this essay, we tackle a collection of such procedures, in four categories, explaining them in a friendly and narrative manner and with an example drawn from the personnel management domain.

La sélection grâce à l'utilisation d'une mesure prédictive : une exploration des méthodes

Louis Laurencelle

Université du Québec à Trois-Rivières

MOTS CLÉS : Sélection, régression linéaire, concomitantes, différentiel de sélection, prédiction du rendement

Il existe peu de textes de référence portant sur le problème de la sélection, en particulier la sélection à partir d'un examen et d'une mesure prédictive. Que ce soit pour le recrutement ou la réaffectation de personnel, l'admission scolaire, la qualification sportive, voire le contrôle de qualité d'un produit usiné, diverses procédures existent, dont les fondements mathématiques restent peu ou mal connus. Nous en abordons ici une panoplie, en les regroupant en quatre catégories, sur un mode convivial et narratif et au gré d'un exemple inspiré de la gestion de personnel.

KEY WORDS : Selection, linear regression, concomitants, selection differential, work output

There are but few textbooks that deal with selection and selection procedures, particularly with the use of a predictive measure. Be it for the recruiting or reassignment of personnel, admission into educational programmes, qualification in sports, even quality control for factory-made products, many procedures are available whose mathematical bases are little known or publicized. In this essay, we tackle a collection of such procedures, in four categories, explaining them in a friendly and narrative manner and with an example drawn from the personnel management domain.

PALAVRAS-CHAVE : Seleccção, regressão linear, concomitantes, diferencial de selecção, predição do rendimento

Existem poucos textos de referência que tratem o problema da selecção, em particular a selecção a partir de um exame e de uma medida preditiva. Seja para o recrutamento ou para a reafecção de pessoal, para a admissão escolar, para a qualificação desportiva, isto é, para o controle de qualidade de um produto fabricado, existem diversos procedimentos, cujos fundamentos matemáticos continuam pouco ou mal conhecidos. Abordaremos aqui uma panóplia deles, reagrupando-os em quatro categorias, de um modo narrativo e a partir de um exemplo inspirado na gestão de pessoal.

Note de l'auteur : Toute correspondance peut être adressée comme suit : Louis Laurencelle, Ph.D., professeur, Département des sciences de l'activité phisique, Université du Québec à Trois-Rivières; téléphone 819-376-5011 (poste 3794); télécopieur 819-376-5092; courriel: Louis_Laurencelle@UQTR.CA

1. Introduction et contextes

La sélection en vue de l'admission dans un programme, de l'engagement dans un poste, de la qualification dans une fonction, est restée un domaine relativement peu défriché par les auteurs, notamment sur le plan des procédés statistiques encourus. Certes, les agences de personnel, les écoles, les grandes entreprises *font de la sélection* et elles le font consciencieusement, de bonne foi. Parmi les outils les plus exploités se trouvent l'entrevue, le CV ou bulletin de qualifications, parfois une épreuve chiffrée ou un test d'aptitude. Même alors, il est rare que la relation existante entre le test d'aptitude et la performance attendue (à l'école ou dans un emploi) soit explicitement connue, et encore plus rares sont les cas où des procédures d'aide à la décision mathématiquement fondées soient disponibles. C'est dans ce contexte que nous proposons au lecteur d'examiner ici quelques procédures possibles, la plupart peu connues, pour aborder scientifiquement le problème de la sélection.

Le champ d'application où l'on retrouve le plus de documentation sur les questions reliées à la sélection est sans contredit celui de la gestion du personnel, dans les entreprises et les grandes organisations; nous entendons par là les grandes usines et manufactures, les bureaux de cols blancs (téléphonie, compagnies d'argent ou de crédit, etc.), la fonction publique, l'armée, ainsi de suite. À côté des grands nombres que suggèrent ces contextes, existent aussi la chasse aux têtes et la sélection du personnel hautement qualifié, qui pour n'impliquer parfois qu'un seul poste à combler n'en posent pas moins de grandes exigences. Les ouvrages de Guion (1965, 1998) abordent tous les aspects de ce problème et sont typiques du genre, comme en français le volume de Pettersen (2000), avec très peu de matière quant aux modes et procédures statistiques de la sélection. Robertson et Smith (1989) de même que, par exemple, Gatewood et Feild (1990) y contribuent plus généreusement. L'ouvrage de référence, qui reste irremplaçable, est sans doute le *Psychological tests and personnel decisions* de Cronbach et Gleser (1965).

Nous avons recensé peu de publications spécifiques sur la sélection en éducation, particulièrement sur les procédures de sélection par *testing*. Cronbach et Gleser (1965) y réfèrent à l'occasion. Hills (1971) explore différentes questions reliées à la sélection, l'admission et le classement des élèves sur un mode narratif, en fournissant une abondante bibliographie. Finney (1962) par ailleurs prolonge Cronbach et Gleser en explorant d'autres modes de sélection, y compris les bases mathématiques sous-jacentes.

Nous envisageons aussi un contexte tout différent, à savoir celui du *contrôle de qualité* en matière de produits usinés, contexte dans lequel une procédure de *testing* permettrait de certifier un produit pour une application particulière et exigeante. L'opération de test et de sélection est d'usage courant en micro-informatique (on teste la vitesse d'opération maximale possible d'un composant, sa résistance à la chaleur, etc.) et ailleurs, les procédures présentées ici pouvant y être de quelque utilité.

Comme nous nous intéressons, dans cet article, à l'appareil mathématique qu'on peut mettre en œuvre dans une démarche de sélection, notre documentation concerne principalement cet aspect. L'univers de référence dans lequel nous nous plaçons est classique : celui de la distribution normale et de la régression linéaire des moindres carrés (Draper & Smith, 1981). En cours d'exposé, nous aborderons des territoires moins connus, d'aucuns tout neufs, de la théorie statistique : nous indiquerons alors nos sources.

Finalement, pour faciliter l'exposé et la continuité des contenus, il était nécessaire pour nous de choisir un champ d'application unique, qui servirait aussi de modèle à adapter aux autres champs. À l'instar des grands auteurs cités, dont Cronbach et Gleser (1965), nous adoptons le champ de la gestion et de l'embauche de personnel, en nous référant à « l'organisation », au « rendement » de ses employés, à ses « candidats » à l'emploi, etc., de manière directe (pour la gestion) ou métaphorique (pour l'éducation, le contrôle de qualité, etc.).

2. Relever le niveau de rendement moyen d'une organisation

Dans une grande organisation, une dans laquelle l'objectif se mesure en termes de rendement global, l'accent porte moins sur le rendement de chaque employé, et la sélection vise plutôt un rendement qui soit satisfaisant *en moyenne*. Il faut tout de même sélectionner les candidats et, pour ce faire, trouver et appliquer un critère qui promette le rendement moyen attendu. Nous avons affaire ici à la technique du *différentiel de sélection* (Brogden, 1946 ; Cronbach & Gleser, 1965 ; Laurencelle, 1998). La décision de sélection peut se présenter sous trois formes, que nous examinons à tour de rôle.

2.1 Nous connaissons le rendement (Y) de chaque candidat, et nous voulons retenir seulement ceux dont le groupe fournira un rendement moyen satisfaisant (μ_Y^*)

Supposons une organisation dans laquelle le rendement des employés a pour mesure Y , et que la moyenne et l'écart-type courants y sont de $\mu_Y = 100$ et $\sigma_Y = 20$. L'employeur veut hausser le rendement moyen à la valeur $\mu_Y^* = 110$, par une sélection interne (assortie de licenciements) ou par des réaffectations. La figure 1 illustre la situation. Quel seuil de sélection Y_C faut-il appliquer ici?

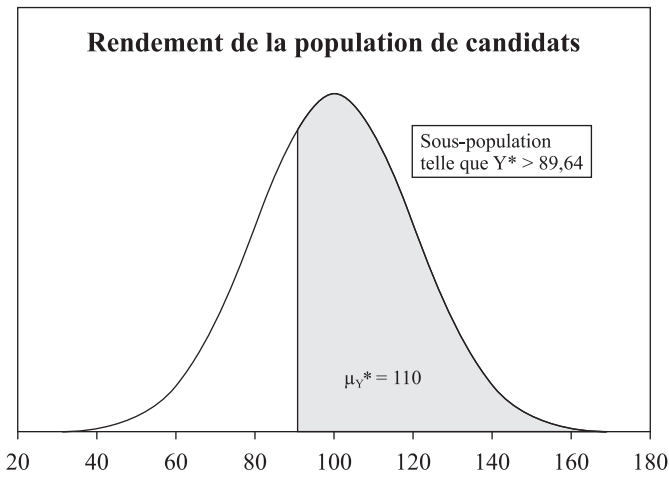


Figure 1. *Distribution de la mesure de rendement (Y) dans la population, selon un modèle normal de moyenne $\mu_Y = 100$ et d'écart-type $\sigma_Y = 20$. La mesure de rendement dans la sous-population est dénotée Y^* .*

Nous connaissons donc le rendement Y de chaque employé ainsi que les statistiques sommaires de l'organisation, μ_Y et σ_Y ; nous adoptons tout au long le modèle de la distribution normale. Comme le montre la figure 1, il s'agit pour nous de trouver un seuil Y_C tel que la moyenne des personnes retenues (*i.e.* pour qui $Y \geq Y_C$) égale la moyenne souhaitée, μ_Y^* . Standardisons d'abord ces quantités,

$$\frac{\mu_Y^* - \mu_Y}{\sigma_Y} = \Delta_C, \quad (1)$$

pour obtenir le différentiel de sélection, Δ_C , soit l'espérance (ou moyenne) des valeurs situées au-dessus d'un seuil C dans une distribution normale standard. Dans une telle distribution, de moyenne 0 et d'écart-type 1, la valeur de Δ_C est déterminée par l'équivalence

$$\Delta_C = \frac{\varphi(C)}{1 - \Phi(C)}, \quad (2)$$

où φ et Φ désignent respectivement la densité de probabilité et la fonction de répartition de la distribution normale standard, évaluées à C . Par exemple, à $C = 1$ ($= z$), $\varphi(C) = \exp(-1/2C^2)/\sqrt{2\pi} = 0,242$, $\Phi(C) = 0,841$ et $\Delta = 1,522$. Le tableau 1 illustre la relation entre Δ_C , d'une part, et le seuil C et la portion d'éléments sélectionnés, $p = 1 - \Phi(C)$ d'autre part.

Pour notre exemple, nous calculons $\Delta = (110-100)/20 = 0,50$; le différentiel $\Delta = 0,50$ correspond à $C = -0,518$ en valeur standardisée. Le score-seuil à appliquer pour la sélection sera donc $Y_C = \mu_Y + \sigma_Y \times C = 100 + 20 \times (-0,518) = 89,64$, correspondant à un taux de sélection $p = 0,698$. Ainsi, en retenant seulement les employés dont le rendement atteint environ $Y \approx 90$ ou mieux, l'employeur espère obtenir un rendement moyen de 110, le pourcentage de rejet étant de $100(1-p) \approx 30\%$.

Tableau 1
*Valeurs associées du différentiel de sélection standard (Δ),
du score-seuil (C) et de la fraction d'éléments sélectionnés (p)*

Δ	C	p	Δ	C	p	Δ	C	p
0,0	$-\infty$	1,000	1,0	0,303	0,381	2,0	1,572	0,058
0,1	-1,691	0,995	1,1	0,443	0,329	2,1	1,688	0,046
0,2	-1,264	0,897	1,2	0,580	0,281	2,2	1,803	0,036
0,3	-0,967	0,833	1,3	0,712	0,238	2,3	1,917	0,028
0,4	-0,726	0,766	1,4	0,842	0,200	2,4	2,030	0,021
0,5	-0,518	0,698	1,5	0,969	0,166	2,5	2,142	0,016
0,6	-0,331	0,630	1,6	1,093	0,137	2,6	2,254	0,012
0,7	-0,158	0,563	1,7	1,215	0,112	2,7	2,365	0,009
0,8	0,003	0,499	1,8	1,336	0,091	2,8	2,475	0,007
0,9	0,156	0,438	1,9	1,454	0,073	2,9	2,585	0,005

2.2 Nous obtenons le score (X) de chaque candidat à un test de sélection, le test ayant une corrélation $\rho_{X,Y}$ connue avec le rendement (Y), et nous voulons retenir ceux qui amèneront le rendement à une valeur moyenne relevée (μ_Y^*).

Le bureau du personnel possède un test d'aptitude au travail, de mesure X, avec moyenne μ_X et écart-type σ_X , lequel test a une corrélation $\rho_{X,Y}$ avec la mesure de rendement. Supposons ici $\mu_X = 60$, $\sigma_X = 10$ et $\rho_{X,Y} = 0,70$. Le rendement typique, ou moyen, des candidats se distribue encore normalement, avec $\mu_Y = 100$ et $\sigma_Y = 20$: quel seuil doit-on utiliser au test de sélection afin d'espérer un rendement moyen de $\mu_Y^* = 110$, de la part des nouveaux candidats?

Si elle existe, la relation linéaire entre prédicteur (X) et critère (Y) a comme forme :

$$\hat{Y} = bX + a, \quad (3)$$

où $b = \rho_{X,Y} \sigma_Y / \sigma_X = 0,70 \times 20 / 10 = 1,4$ et $a = \mu_Y - b \mu_X = 100 - 1,4 \times 60 = 16$. La rétention de tous les candidats (toutes les valeurs de X) aboutit à l'obtention d'un rendement moyen, près de $\mu_Y = 100$; il nous faut donc retenir seulement les candidats excédant un certain score-seuil en X, disons X_C , afin d'assurer la moyenne haussée de $\mu_Y^* = 110$. L'égalité à observer est :

$$\mu_Y^* = b \mu_X^* + a; \quad (4)$$

après un peu d'algèbre, la quantité déterminante est le différentiel $\Delta_X = (\mu_Y^* - \mu_Y) / (\rho \sigma_Y)$, qui vaut ici $(110 - 100) / (0,7 \times 20) \approx 0,714$. Au tableau 1, par interpolation sur Δ entre 0,7 et 0,8, on trouve, pour $\Delta = 0,714$, $C \approx -0,135$ et $p \approx 0,554$. Le score-seuil à appliquer au test pour recruter un candidat sera donc $X_C = \mu_X + C \sigma_X = 60 + (-0,135) \times 10 \approx 58,65$, le pourcentage de rejet des candidats étant ici de $100(1-p) \approx 45\%$.

2.3 En plaçant le candidat en situation de travail, nous obtenons une mesure de rendement estimative (Y'), dont la fidélité est $\rho_{YY'}$, et voulons retenir ceux dont le rendement réel (Y) gravitera vers une moyenne relevée (μ_Y^*).

Au lieu d'un test prédictif (X) ou du rendement réel (Y) de la personne, nous voici en présence d'une estimation de son rendement (Y'), obtenue par exemple en une demi-journée ou une journée de travail; cette estimation étant de courte durée, sa fidélité $\rho_{YY'}$ reste imparfaite, disons égale à 0,81. Comment procéder?

Le rendement étant le rendement, on peut supposer que la moyenne ($\mu_{Y'}$) reste la même ($= \mu_Y = 100$). L'écart-type, cependant, reflète maintenant une composante d'erreur de mesure et devient plus grand, soit $\sigma_{Y'} = \sigma_Y / \sqrt{\rho_{YY'}}$ ($\approx 22,22$). Le reste de la situation peut se traiter comme si le rendement estimatif était un prédicteur (voir par. 2.2). Le différentiel à trouver ici est donné par $\Delta = (\mu_Y^* - \mu_Y) / (\rho_{XX} \sigma_{Y'}) = (110 - 100) / (0,81 \times 22,22) \approx 0,556$. De là, par le tableau 1, $C \approx -0,413$ et $p = 0,660$. Le seuil de sélection à exiger, $Y' \geq Y_C$, est donné par $\mu_{Y'} + C\sigma_{Y'} \approx 100 + (-0,413) \times 22,22 \approx 90,82$, la procédure entraînant un taux de rejet de 34%.

3. Façonner un critère de sélection prédictive tel que les candidats retenus soient assurés d'atteindre un rendement cible

Un candidat obtient, au test d'aptitude, le score X_o : quelle probabilité avons-nous que ce candidat atteigne le niveau de rendement demandé, Y_C ? Aussi, peut-on mettre sur pied un seuil de sélection (X_C) tel que chaque candidat admis ait une probabilité d'au moins β d'atteindre le rendement demandé Y_C ?

3.1 Robert obtient le score $X_o = 68$ au test d'aptitude : quelle est sa probabilité de fournir éventuellement le rendement $Y_C = 110$ demandé? Quelle valeur de seuil doit-on exiger pour que chaque candidat retenu soit assuré de fournir un tel rendement?

Dans un système à régression linéaire tel que (3) et en admettant les conditions habituelles de normalité et d'homoscédasticité, le comportement de Y conditionné par $X = X_o$ suit simplement une distribution normale, avec pour moyenne $bX_o + a$ et pour variance $\sigma_Y^2(1-\rho^2)$. Ainsi, pour notre candidat qui obtient $X_o = 68$, on peut établir ses chances d'être satisfaisant à l'emploi par :

$$\Pr\{ Z \geq (Y_C - b X_o - a) / \sigma_Y \sqrt{1 - \rho^2} \} ; \quad (5)$$

pour Robert, en utilisant $b = 1,4$, $a = 16$, $\sigma_Y = 20$ et $\rho = 0,70$, nous calculons $\Pr\{ Z \geq (110 - 1,4 \times 68 - 16) / 20 \sqrt{1 - 0,70^2} \} \approx \Pr\{ Z \geq -0,084 \} \approx 0,533$, la probabilité équivalant à une chance sur deux.

Si dans cette population de candidats on exige de chacun une garantie d'au moins $\beta = 0,90$ à l'effet qu'il fournisse un rendement satisfaisant, la valeur X_o^* à appliquer au test d'admission peut s'obtenir en inversant l'expression (5) et elle est donnée par :

$$X_o^* = [Y_C - a + \Phi^{-1}(\beta) \sigma_Y \sqrt{1 - \rho^2}] / b ; \quad (6)$$

en utilisant $\Phi^{-1}(\beta) = \Phi^{-1}(0,90) = 1,282$, soit le 90^e centile de la distribution normale standard, nous obtenons $X_o^* = [110 - 16 + 1,282 \times 20 \times \sqrt{1 - 0,70^2}] / 1,4 \approx 80,22$, un seuil assez sévère.

3.2 Nous obtenons le score (X) de chaque candidat à un test de sélection, et nous souhaitons recruter tous ceux dont le rendement sera satisfaisant avec une assurance moyenne de β .

Au lieu de l'ensemble des candidats postulant l'emploi, il est possible de considérer seulement ceux qui, au test de sélection, obtiendront un score minimal X_C ou mieux. Ces éléments, définis par $X \geq X_C$, forment une *sous-population* dont la probabilité de haut rendement est plus élevée. En fait, le rendement Y^* dans cette sous-population dépend d'un mélange de distributions, exprimé en unités standards par la densité :

$$g(Z_Y^* | C) = \int_C^\infty \frac{\varphi(u)}{1 - \Phi(C)} \frac{\varphi[(Y - \rho u) / \sqrt{1 - \rho^2}]}{\sqrt{1 - \rho^2}} du \quad (7)$$

Le seuil d'admission X_C pour cette sous-population influence directement la distribution des rendements et, par conséquent, détermine la probabilité que n'importe quel membre de cette sous-population fournisse le rendement exigé. Le niveau de rendement demandé, $Y_C = 110$, devient en unités standards $Z_{Y_C} = (110 - 100) / 20 = 0,50$. Pour que $\Pr\{Z_Y \geq 0,50 | X_C\} = \beta = 0,90$ avec $\rho = 0,70$, l'intégration de (7) montre qu'on doit fixer $C = 1,683$, d'où en unités directes $X_C = 60 + 10 \times 1,683 \approx 76,83$. Ainsi, en admettant tous les candidats qui, au test de sélection, obtiennent 76,83 ou mieux, on est assuré à 90% ($= 100\beta$) que tous donneront un rendement d'au moins $Y_C = 110$ ¹. Laurencelle (2005) documente la distribution (7).

4. Façonner un critère de sélection prédictive tel que le plus performant (le moins performant) parmi les k meilleurs candidats retenus soit assuré d'atteindre un rendement cible

Parfois, la sélection s'effectue par cohorte : l'organisation ouvre ses portes pour examiner N candidats, p. ex. $N = 10, 25$ ou 200 , et elle espère recruter parmi eux k candidats valables grâce à un examen de sélection. Pour autant qu'il faille engager un nombre k de candidats, il y a lieu de s'interroger sur les chances que tous les candidats retenus soient capables du rendement demandé, ou sur celles qu'au moins le meilleur d'entre eux en soit capable.

Reprenons l'exemple courant, dans lequel le rendement Y ($\mu_Y = 100, \sigma_Y = 20$) est associé au prédicteur X ($\mu_X = 60, \sigma_X = 10$) par une relation linéaire de corrélation $\rho = 0,70$, les deux quantités X et Y étant distribuées normalement. Définissons un échantillon de N candidats, évalués en X , soit la série $\{X_1, X_2, \dots, X_N\}$, dont nous retenons ceux obtenant les k meilleurs scores, soit :

$$(X_{(N-k+1)} \ X_{(N-k+2)} \ \dots \ X_{(N)}) ;$$

il s'agit là des *statistiques d'ordre* de la série des N résultats, en fait les k plus hautes statistiques d'ordre. Par exemple, $X_{(N)}$ dénote le meilleur des N scores obtenus en X . À ces statistiques d'ordre $X_{(i)}$, sont associées une à une les *concomitantes* $Y_{[i]}$, l'association observant encore la régression linéaire, soit :

$$Y_{[i]} = \rho X_{(i)} + \sqrt{1-\rho^2} Z, \quad (8)$$

en unités standards de X , Y et Z , cette dernière quantité Z étant une composante aléatoire normale standard. David et Nagaraja (1998) et Yang (1977) développent la théorie mathématique des concomitantes. Or, même si $X_{(N)}$ dénote le meilleur score en X parmi les N candidats, rien n'assure que sa valeur associée $Y_{[N]}$ indique le meilleur rendement: de fait, en l'absence de relation linéaire (si $\rho = 0$), il n'y aurait nul lien entre X et Y .

Choisir «le plus performant» parmi les k meilleurs candidats revient à s'intéresser au maximum des k concomitantes supérieures, soit:

$$V_{k,N} = \max \{ Y_{[N-k+1]} Y_{[N-k+2]} \dots Y_{[N]} \} \quad (9)$$

et à sa distribution. Si, par contre, ce sont tous les k meilleurs candidats dont on espère un rendement donné, le rendement du candidat le moins performant devient la variable clé, et la statistique appropriée est maintenant le minimum des k concomitantes supérieures, soit:

$$W_{k,N} = \min \{ Y_{[N-k+1]} Y_{[N-k+2]} \dots Y_{[N]} \} . \quad (10)$$

Nagaraja et David (1994) proposent une expression pour la fonction de répartition de $V_{k,N}$ et des approximations, notamment pour l'espérance mathématique; Laurencelle (2005) propose un algorithme de calcul pour les densités de $V_{k,N}$ et $W_{k,N}$ ainsi que des estimations Monte Carlo. Notons que, lorsque $k = 1$, nous avons $V_{k,N} = W_{k,N} = Y_{[N]}$, les quantités coïncidant alors avec la concomitante la plus haute.

4.1 Nous avons $N = 10$ candidats; quelle est la probabilité que le «meilleur», celui qui obtient le meilleur score au test de sélection, atteigne le rendement demandé?

Nous nous intéressons ici à seulement un candidat parmi les dix examinés, soit le meilleur en X , selon $k = 1$ et $N = 10$. La distribution de $X_{(10)}$, plus spécifiquement $X_{(10:10)}$, dans la population est connue, de même que celle de $Y_{[10:10]}$ (Laurencelle, 2005). En unités standards, le rendement à atteindre ($Y_C = 110$) vaut $(110-100)/20 = 0,50$. Utilisant aussi $\rho = 0,70$, l'évaluation numérique permet d'obtenir, en unités standards:

$$\Pr \{ V_{1,10} \geq 0,50 \} = \Pr \{ Y_{(10:10)} \geq 0,50 \} \approx 0,76 .$$

L'assurance de rendement voulue, à $\beta = 0,90$, n'est pas atteinte.

4.2 Dans le contexte ci-dessus, comment l'organisation peut-elle s'assurer de retenir un candidat satisfaisant ?

En retenant les $k = 2$ candidats présentant les meilleurs scores à l'examen de sélection, nous améliorons nos chances de recruter *le* candidat qui fournira *le meilleur rendement*, donc la probabilité d'atteindre, avec lui, le rendement demandé. En utilisant une cohorte de $N = 10$ candidats et en retenant les $k = 2$ meilleurs, nous avons maintenant affaire au maximum des deux plus hautes concomitantes, $V_{2,10}$, et, dans le contexte connu, l'estimation Monte Carlo nous donne :

$$\Pr\{V_{2,10} \geq 0,50\} \approx 0,89,$$

atteignant presque $\beta = 0,90$. En fait, avec $k = 3$, nous obtiendrions $\Pr\{V_{3,10} \geq 0,50\} \approx 0,93$.

Il existe au moins deux autres moyens pour atteindre l'assurance β requise, que ce soit pour l'engagement d'un seul ou de plusieurs candidats satisfaisants. Ces deux moyens sont exposés ci-après.

4.3 Comment assurer autrement que l'un des deux candidats retenus soit satisfaisant ?

La méthode la plus simple, avec $k = 2$ candidats retenus (et engagés à titre probatoire dans l'organisation), est d'augmenter la cohorte de candidats. En fait, avec $N = 10$, la probabilité obtenue pour $V_{2,10}$ est d'environ 0,887 ; avec $N = 11$, elle monte à 0,898, et elle atteint 0,908 avec $N = 12$. Ainsi, en engageant les deux meilleurs candidats parmi 12, l'employeur s'assure de compter dans ses rangs au moins un candidat satisfaisant.

Une autre méthode consiste à restreindre la sélection à une sous-population, telle que les candidats excèdent tous un score-seuil, $X \geq X_C$, plutôt que de considérer les candidats de toute valeur : cela suppose une présélection. Dans le cas présent, où $\rho = 0,70$ et où l'on exige une assurance de $\beta = 0,90$ que le meilleur des deux candidats retenus donne un rendement Y_C équivalant à 0,50 en unités standards, il faudrait restreindre la sélection aux candidats pour qui $Z_X \geq -1,33$, soit $X \geq 60 + (-1,33) \times 10 \approx 46,70$. Nous écririons :

$$\Pr\{V_{2,10} \mid Z_X \geq -1,33\} \approx 0,90 ;$$

en d'autres mots, sur dix candidats satisfaisant au critère $Z_C \geq -1,33$ ou $X \geq 46,70$ et en retenant ceux qui produisent les deux meilleurs scores, on est assuré à $\beta = 0,90$ que l'un des deux au moins fournira le rendement demandé.

4.4 Comment assurer que, sur les deux candidats retenus (parmi dix), les deux aient un rendement satisfaisant ?

Pour assurer, à partir de N candidats, que les k candidats retenus soient satisfaisants, il faut que le moins performant de ceux-ci atteigne le rendement demandé : un critère plus exigeant, exprimé généralement par :

$$\Pr \{ W_{k,N} \geq Z_{Yc} \} = \beta .$$

Ainsi, sur $N = 10$ candidats examinés sans restriction, la probabilité que les $k = 2$ meilleurs en X donnent le rendement attendu n'est que de 0,47. Pour atteindre $\beta = 0,90$, il faut ou bien examiner $N \approx 76$ candidats de tout acabit, ou bien, avec $N = 10$ candidats, restreindre ceux-ci à une sous-population telle que $Z_X \geq 1,48$ ou $X \geq 74,80$. Noter que toutes les valeurs fournies ici sont issues d'estimations Monte Carlo très précises (utilisant 10^6 échantillons aléatoires).

Le procédé consistant à limiter la sélection à une sous-population n'est pas sans coût : le dernier cas ci-dessus en est un exemple. Dans la population générale (à distribution normale), la fraction d'éléments ayant une mesure supérieure à $Z_X = 1,48$ est d'environ 0,069 : il s'agit donc d'individus peu faciles à trouver. En fait, il faut en moyenne examiner 145 (± 44) éléments pour en trouver dix de cette force. En pratique toutefois, ceux qui se présentent volontairement à un test de sélection ne proviennent guère de la zone inférieure de la population, de sorte que cette exigence est souvent moins astreignante qu'elle ne paraît.

5. Façonner un critère de sélection prédictive tel que les s meilleurs candidats soient assurés d'être retenus

Pour occuper, par exemple, deux emplois dans une organisation, si dix candidats se présentent à la sélection, le choix des candidats obtenant les deux meilleurs scores au test (X) est évidemment le plus judicieux. Mais, ce faisant, est-on assuré d'avoir recruté parmi les candidats disponibles les deux qui auraient fourni le meilleur rendement ? Notons qu'ici l'organisation ne fixe pas un niveau (Y) de rendement demandé : l'employeur souhaite seulement ne pas manquer de recruter les deux bons candidats dont il a besoin, quitte à engager un ou quelques candidats de plus sur une base temporaire.

Le problème de sélection se pose comme suit. Nous avons une population de candidats, chacun étant évaluable par un test d'aptitude à l'emploi, de mesure X , et chacun étant susceptible de fournir un rendement au travail de mesure Y ; la corrélation linéaire entre mesure d'aptitude et mesure de rendement est $\rho_{X,Y}$. À partir d'une cohorte de N candidats, l'organisation recrute (temporairement) r personnes dans le but de retenir éventuellement les s meilleurs au rendement, avec une probabilité ou assurance β ($0 < \beta < 1$). On peut voir le problème comme une probabilité de capture, correspondant symboliquement à :

$$\Pr\{ C(N, r) = s \mid \rho \} = \beta . \quad (11)$$

Par exemple, avec $N = 10$ candidats et si $\rho = 0,70$, on évalue $\Pr\{ C(10, 2) = 2 \mid \rho = 0,70 \} \approx 0,23$. Ainsi, avec cette corrélation entre test et rendement et à partir d'une cohorte de dix candidats, il y a moins d'une chance sur quatre que les deux candidats les plus performants soient aussi ceux obtenant les meilleurs scores au test de sélection, donc peu de chances que, en recrutant les deux meilleurs candidats au test de sélection, l'employeur mette la main sur les deux employés les plus performants.

Les mathématiques reliées à cette question sont, comme à l'accoutumée, relativement simples pour les cas où $\rho = 0,00$ et $\rho = 1,00$ mais d'une difficulté ahurissante autrement (David, O'Connell & Yang, 1977; Laurencelle, 2005; Yeo & David, 1984); le recours à l'estimation Monte Carlo reste incontournable.

5.1 Comment assurer que, sur N candidats, la sélection permette de retenir les s meilleurs (pour ce qui est du rendement) ? au moins l'un des s meilleurs ?

La taille (N) de la cohorte, dans cette procédure, est fixée d'avance, par exemple $N = 10$. Bien sûr, dans le but de recruter des candidats pour ($s =$) 2 emplois, ceux-là auront, en probabilité, un rendement encore meilleur s'ils sont tirés d'une cohorte plus nombreuse, comme $N = 30$.

Revenons à notre cohorte de $N = 10$ candidats et un test de sélection (X) corrélant à $\rho = 0,70$ avec le rendement (Y). Si l'on recrute tous les candidats, *i.e.* $r = N = 10$, on est évidemment assuré à 100% d'en retenir les $s = 2$ meilleurs. Toutefois, le coût encouru par l'emploi de dix candidats afin de retenir les deux meilleurs avec certitude (*i.e.* $\beta = 1,00$) est très élevé, souvent trop élevé. Si l'on réduit notre niveau d'assurance β à une valeur raisonnable, comme $\beta = 0,50$, il suffira alors de recruter $r = 4$ candidats ($\Pr \approx 0,620$), ou peut-être $r \approx 7$ ($\Pr \approx 0,920$) pour atteindre notre niveau $\beta = 0,90$.

D'autres stratégies de recrutement amèneraient d'autres critères. L'employeur, qui ciblerait les s meilleurs candidats, pourrait vouloir en retenir x ($1 \leq x \leq s$) parmi ceux-là en recrutant r candidats. La probabilité correspondante² est alors :

$$\Pr\{C(N, r, s) = x \mid \rho\} = \beta. \quad (12)$$

la quantité (11) équivalant au cas particulier $\Pr\{C(N, r, s) = s \mid \rho\}$. Par exemple, parmi les deux meilleurs candidats, la probabilité (avec $\rho = 0,70$) d'en retenir au moins un si on recrute les deux meilleurs au test est de 0,84, et elle passe à $\Pr\{C(10, 3, 2) = 1\} \approx 0,93$ en recrutant $r = 3$ candidats. Dans la même ligne stratégique et afin d'attribuer $x = 2$ emplois à deux des $s = 3$ meilleurs performeurs, nous obtenons $\Pr\{C(10, 2, 3) = 2\} \approx 0,45$, la probabilité devenant satisfaisante dès $r = 5$, avec $\Pr\{C(10, 5, 3) = 2\} \approx 0,94$.

Le choix de la stratégie de sélection dépend des besoins de l'organisation, de la disponibilité des candidats, du coût de sélection, du coût d'un emploi en phase probatoire, du bénéfice au rendement et d'autres considérations : nous n'en avons effleuré que les aspects statistiques.

6. Considérations générales

6.1 Les populations et les modèles paramétriques

Les modèles de sélection évoqués ci-dessus concernent des candidats et une organisation qui cherche à combler un poste, ou plusieurs postes, en escomptant des recrues un certain niveau de rendement. Pour coucher ces concepts et ces actions dans un modèle opérationnel, il est nécessaire de stipuler une *population de candidats*, une *population d'employés*, de même que des mesures de qualification et de rendement, que nous avons dénotées X et Y , respectivement. À l'instar des autres auteurs, nous invoquons aussi une relation, une régression linéaire entre Y et X ; cette relation *linéaire* reste d'autant plus plausible, selon nous, que X est souvent conçu comme une variable précurseur du rendement Y , voire un rendement Y en miniature, des relations autres que linéaires restant une curiosité.

Cela dit, en quoi consiste vraiment la population de candidats mentionnée ? Cette population coïncide difficilement avec la population générale, puisque ce n'est pas n'importe quels individus qui répondent aux offres d'emploi ou se présentent au bureau du personnel d'une entreprise. C'est pourquoi nous préférons parler d'une *population de candidats*, laquelle

toutefois a peu de chances d'être recensée et mesurée et pour laquelle les statistiques descriptives risquent de faire défaut. Une réserve de même nature s'applique aux employés, dont le rendement (mesurable, par exemple, relativement au bénéfice net par année, par employé) risque de varier selon l'âge, l'expérience et la sécurité d'emploi de chacun. Le modèle normal appliqué à cette population est, encore une fois, une fiction commode, qui recouvre en fait un mélange inextricable de distributions plus ou moins différentes ou mal connues.

6.2 Les valeurs normatives (incluant ρ)

Admettons à toutes fins utiles la réalité heuristique des modèles normaux proposés; il reste à ajuster ces modèles à la réalité particulière d'une entreprise, d'une sélection. Nous ne disposons pour ainsi dire jamais des données de toute une population, particulièrement la population virtuelle des candidats à l'emploi: avons-nous assez de données pour exploiter avec confiance les modèles paramétriques invoqués?

Dans l'exposé de procédures présenté ici, nous nous sommes exprimé par des valeurs paramétriques, soit μ_X et σ_X pour le modèle des scores d'aptitude, μ_Y et σ_Y pour la mesure de rendement, et $\rho_{X,Y}$ pour la valeur de corrélation entre les deux. Si le modèle est connu et normal, ou stipulé normal, les estimateurs habituels de moyenne ($\bar{X}$) et d'écart-type (s_X) sont habituellement satisfaisants, même s'ils sont obtenus par un échantillon normatif de taille (N) modeste. Rappelons que la marge d'erreur typique (ou erreur-type) de $\bar{X}$ est $\sigma_X/\sqrt{N}$, alors que pour l'écart-type, elle est d'environ $\sigma_X/\sqrt{2N}$.

Le talon d'Achille du système prédictif concerne l'estimateur de corrélation, r , qui joue un rôle critique dans les procédures de sélection décrites, puisque c'est sa valeur qui est substituée au coefficient de validité prédictive $\rho_{X,Y}$, au cœur de toutes nos équations. Or, le coefficient r est un estimateur relativement imprécis (Laurencelle, 2000), son erreur-type approchant $1/\sqrt{N}$. Par exemple, pour une valeur paramétrique $\rho = 0,40$, l'estimateur r sur un échantillon de $N = 100$ couples (X,Y) a comme erreur-type 0,0848 et, pour intervalle de confiance à 95 %, l'étendue (0,223 ; 0,554), l'incertitude sur la valeur de ρ étant beaucoup trop forte pour légitimer la sélection équitable des candidats. Par bonheur, Robertson et Smith (1989), qui ont recensé un grand nombre d'applications et d'études sur la sélection de personnels, rapportent avoir trouvé des tailles d'échantillons très élevées, le nombre $N = 3000$ indiquant une hauteur typique. D'autre part, même si nul auteur n'en fait la remarque, les procédures de sélection de type: «prérendement (Y') pour

rendement(Y)», mentionnées au paragraphe 2.3, font appel à un schéma d'évaluation semblable à celui de l'estimation de la fidélité (Laurencelle, 1998; Laveault & Grégoire, 2002). Le coefficient de corrélation appliqué, correspondant à la racine carrée du coefficient de fidélité, $\sqrt{\rho_{YY}}$, est une valeur habituellement très forte ($> 0,80$), dotée par conséquent d'une erreur-type plus faible : ces procédures se révéleraient donc plutôt robustes, même en l'absence d'un échantillon normatif important.

6.3 *Le modèle décisionnel global*

Nous n'avons abordé, dans notre exposé, qu'une petite partie du problème global de la sélection, et nous avons schématisé cette partie dans le modèle d'une distribution bivariée normale. Comme nous l'avons évoqué en introduction, le problème global est très complexe (comme tout problème réel) et il est impossible de le réduire en une seule formule mathématique, fût-elle bien savante.

L'approche la plus répandue pour tenter d'aborder le problème global de la sélection consiste sans doute en l'*analyse utilitaire* («utility analysis») qu'on retrouve, par exemple, dans Cronbach et Gleser (1965), Guion (1998) et Pettersen (2000), et qui traduit les éléments du problème dans un langage de gestion financière, en fonction de coûts (en sélection) et de bénéfices (au rendement) : l'élément clé de cette analyse est la quantité Δ_C , que nous appelons le différentiel de sélection (Laurencelle, 1998). Un modèle plus compréhensif inclurait, par exemple, les salaires et indemnités pour les candidats recrutés sur une base temporaire puis licenciés, et les coûts variables de sélection dans les cas de sélection en deux phases.

Chaque entreprise, chaque système scolaire, chaque industrie est éventuellement tenue à élaborer son propre modèle.

6.4 *Les infrastructures mathématiques*

Du point de vue de leur solution mathématique explicite, la plupart des procédures examinées ici posent un défi quasi insurmontable. Au bout facile du spectre, on trouve la valeur Y associée par régression linéaire à une valeur X_0 donnée (voir par. 3.1) : une simple distribution normale, à paramètres déterminés, donne la réponse. À l'autre bout, on a affaire non seulement à une concomitante reliée à une statistique d'ordre, mais surtout aux rangs d'un sous-ensemble de concomitantes et à l'évaluation de probabilités de capture, tel que dans les expressions (11) et (12). Pour ce cas, comme aussi pour d'autres procédures, une solution mathématique explicite n'est possible et

pratique que sous des conditions ultra-simplifiées, peu représentatives de la réalité. Par bonheur, la vitesse des ordinateurs récents et l'utilisation d'algorithmes de calcul efficaces rendent possibles des estimations Monte Carlo de grande précision, basées par exemple sur 10^6 itérations ou davantage.

Laurencelle (2005) inventorie les différents systèmes mathématiques qui fondent les procédures de sélection présentées dans cet article, et il en indique les solutions pratiques.

6.5 En guise de conclusion

Les méthodes systématiques de sélection individuelle, conçues pour recruter des individus qui vont satisfaire un critère de rendement prédéfini, sont nées, ont évolué et continuent de progresser avec les méthodes statistiques qui leur sont contemporaines. À notre avis, il serait difficile de déterminer lequel des deux champs a, historiquement, servi d'inspiration à l'autre. De nouveaux concepts ont vu le jour avec les nouvelles méthodes, et le calcul s'est raffiné et complexifié, allant parfois jusqu'à nécessiter le recours aux estimations Monte Carlo.

Nouvelles méthodes et nouveaux concepts ont permis l'élaboration de nouveaux critères, des critères par exemple qui garantissent (avec un degré de probabilité) le rendement particulier d'un candidat, ou le rendement moyen d'une catégorie de candidats, voire leur rendement minimal. Les organisations qui sont soucieuses de performance – les centres de formation, les industries, les usines de production en série – devront mettre en balance le bénéfice escompté et le coût de mise en œuvre afin de décider de l'utilisation de ces nouveaux outils. Qu'elles soient ou non d'application réaliste dans un avenir proche, il reste que ces méthodes posent des défis fascinants à la fois sur le plan des définitions conceptuelles, des procédures de sélection et des techniques de calcul : une mine d'or pour les chercheurs d'aujourd'hui.

NOTES

1. Le rendement moyen espéré de ces candidats (μ_Y^*) est fourni par un calcul semblable à celui du point (2.2) ci-dessus. Le différentiel attaché à $C = 1,683$ est $\Delta \approx 2,096$, d'où $\mu_Y^* \approx 100 + \sigma_Y \rho_{X,Y} \Delta = 100 + 20 \times 0,70 \times 1,683 \approx 129,34$.
2. Si X et Y sont indépendants avec $\rho = 0$, la probabilité est $\Pr\{C(N, r, s) = x\} = \frac{N-s}{N} C_{r-x} \times \frac{s}{N} C_x / N C_r$.

RÉFÉRENCES

- Cronbach, L.J., & Gleser, G. (1965). *Psychological tests and personnel decisions*. Chicago: University of Chicago Press.
- Brogden, H.E. (1946). On the interpretation of the correlation coefficient as a measure of predictive efficiency. *Journal of educational psychology*, 37, 65-76.
- David, H.A., & Nagaraja, H.N. (1998). Concomitants of order statistics. In N. Balakrishnan & C.R. Rao (dir.), *Handbook of statistics, volume 18. Order statistics: Theory and methods* (pp. 487-513). Elsevier-Science.
- David, H.A., O'Connell, M.J., & Yang, S.S. (1977). Distribution and expected value of the rank of a concomitant of an order statistic. *Annals of statistics*, 5, 216-223.
- Draper, N., & Smith, H. (1981). *Applied regression analysis* (2^e édition). New York: Wiley.
- Finney, D.J. (1962). The statistical evaluation of educational allocation and selection. *Journal of the Royal Statistical Society (series A)*, 125, 525-564.
- Gatewood, R.D., & Field, H.S. (1990). *Human resource selection* (2^e édition). The Dryden Press.
- Guion, R.M. (1965). *Personnel testing*. New York: McGraw-Hill.
- Guion, R.M. (1998). *Assessment, measurement, and prediction for personnel decisions*. Mahwah (NJ): Lawrence Erlbaum Associates.
- Hills, J.R. (1971). Use of measurement in selection and placement. In R.L. Thorndike (dir.), *Educational measurement* (2^e édition) (pp. 680-732). Washington, D.C.: American Council on Education.
- Laurençelle, L. (1998). *Théorie et techniques de la mesure instrumentale*. Sainte-Foy: Presses de l'Université du Québec.
- Laurençelle, L. (2000). Distribution, moments et estimation de la corrélation dans une population normale bivariable. *Lettres statistiques*, 11, 31-58.
- Laurençelle, L. (2005). Les concomitantes et leurs applications possibles pour la sélection. *Lettres statistiques*, 12.
- Laveault, D., & Grégoire, J. (2002). *Introduction aux théories des tests en psychologie et en sciences de l'éducation*. Bruxelles: De Boeck.
- Nagaraja, H.N., & David, H.A. (1994). Distribution of the maximum of concomitants of selected order statistics. *Annals of statistics*, 22, 478-494.
- Pettersen, N. (2000). *Évaluation du potentiel humain dans les organisations*. Sainte-Foy: Presses de l'Université du Québec.
- Robertson, I.T., & Smith, M. (1989). Personnel selection methods. In M. Smith & I.T. Robertson (dir.), *Advances in selection and assessment* (pp. 89-112). New York: Wiley.
- Yang, S.S. (1977). General distribution theory of the concomitants of order statistics. *Annals of statistics*, 5, 996-1002.
- Yeo, W.B., & David, H.A. (1984). Selection through an associated characteristic, with applications to the random effects model. *Journal of the American statistical association*, 79, 399-405.