

Estimation de la densité d'ours noirs avec la
technique de capture-marquage-recapture
par génotypage des poils : revue de la
littérature, résultats des inventaires réalisés
en Abibi-Témiscamingue de 2001 à 2003 et
recommandations



Estimation de la densité d'ours noirs avec la
technique de capture-marquage-recapture par
génotypage des poils : revue de la littérature,
résultats des inventaires réalisés en Abibi-
Témiscamingue de 2001 à 2003 et
recommandations

Novembre 2014

Direction de la faune terrestre et de l'avifaune
Direction générale de l'expertise sur la faune et ses habitats
Secteur de la faune et des parcs

**Forêts, Faune
et Parcs**

Québec 

ÉQUIPE DE RÉALISATION

Auteurs : Sabrina Plante
Contractuelle

Christian Dussault
Ministère des Forêts, de la Faune et des Parcs

Sophie Massé
Ministère des Forêts, de la Faune et des Parcs

Sébastien Lefort
Ministère des Forêts, de la Faune et des Parcs

Référence à citer :

PLANTE, S., C. Dussault, S. Massé et S. Lefort. 2014. *Estimation de la densité d'ours noirs avec la technique de capture-marquage-recapture par génotypage des poils : revue de la littérature, résultats des inventaires réalisés en Abitibi-Témiscamingue de 2001 à 2003 et recommandations*. Québec : ministère des Forêts, de la Faune et des Parcs, Direction générale de l'expertise sur la faune et ses habitats, 136 p.

Dépôt légal – Bibliothèque et Archives nationales du Québec, 2014

ISBN : 978-2-550-71799-7 (version imprimée)

978-2-550-71800-0 (version PDF)

RÉSUMÉ

La qualité de la gestion d'une espèce repose, entre autres, sur nos connaissances de son abondance et de sa dynamique de population. Alors que les populations de cervidés font l'objet d'inventaires aériens en hiver, il n'existe actuellement pas de méthode d'inventaire établie pour l'ours noir au Québec. La gestion de l'ours noir repose donc sur les simulations de populations. Cette réalité devient problématique, car une partie des citoyens associe trop rapidement les cas d'ours importuns à une hausse des densités de population d'ours. Le Ministère souhaite donc se doter d'une méthode d'estimation de la densité d'ours noirs afin d'améliorer la gestion de cette ressource. Le but du présent rapport était d'évaluer la faisabilité et l'efficacité de la méthode de capture-marquage-recapture (CMR) utilisant la reconnaissance individuelle par le génotypage des poils pour l'estimation de la densité d'ours noirs au Québec. Il s'agit d'une méthode non invasive, c'est-à-dire qui ne nécessite pas la capture des individus, et qui a été utilisée avec succès ailleurs dans le monde. Au Québec, cette méthode a également fait l'objet d'une expérimentation en Abitibi-Témiscamingue entre 2001 et 2003. Nous avons d'abord réalisé une revue de littérature sur les techniques actuellement utilisées pour faire l'inventaire des ursidés en mettant l'accent sur la méthode de CMR par génotypage des poils. Nous avons ensuite analysé les résultats des projets expérimentaux réalisés en Abitibi-Témiscamingue. Suite à ces analyses et à l'ensemble de notre démarche, nous considérons que cette méthode serait applicable et avantageuse dans le cadre du suivi des populations d'ours noirs au Québec. En nous basant sur nos lectures et sur les expériences passées, nous proposons une méthode d'inventaire des populations d'ours noirs adaptée au contexte québécois.

TABLE DES MATIÈRES

RÉSUMÉ	iv
TABLE DES MATIÈRES	v
LISTE DES TABLEAUX	viii
LISTE DES FIGURES	x
LISTE DES ANNEXES	xi
1. INTRODUCTION	1
2. OBJECTIFS	4
3. REVUE DE LITTÉRATURE	5
3.1. Techniques d’inventaire de l’ours noir utilisant une approche de CMR.....	5
<i>Description de la méthode de CMR avec identification individuelle par génotypage ...</i>	<i>5</i>
<i>Caractéristiques de l’aire d’étude et disposition des blocs pour un inventaire de CMR par génotypage</i>	<i>8</i>
<i>Nombre et disposition des stations d’échantillonnage</i>	<i>9</i>
<i>Pièges à poils, appâts et leurres olfactifs aux stations d’échantillonnage.....</i>	<i>11</i>
<i>Période et durée de l’échantillonnage</i>	<i>13</i>
<i>Récolte et conservation des échantillons de poils</i>	<i>15</i>
<i>Analyse génétique des poils.....</i>	<i>16</i>
<i>Sexage des individus</i>	<i>22</i>
<i>Identification des génotypes identiques</i>	<i>23</i>
<i>Estimation de la taille de la population (N)</i>	<i>23</i>
<i>Estimation de la densité d’une population</i>	<i>27</i>
<i>Sources de biais des méthodes de CMR par génotypage</i>	<i>33</i>
<i>Sources d’erreurs des analyses génétiques et solutions</i>	<i>37</i>
4. INVENTAIRES RÉALISÉS EN ABITIBI-TÉMISCAMINGUE (2001-2003).....	44
4.1. Description des études	44

<i>Mise en contexte</i>	44
<i>Sites d'étude</i>	45
<i>Méthodes de récolte des poils</i>	46
4.2. Analyses génétiques et démographiques	47
<i>Analyses génétiques</i>	47
<i>Identification des génotypes identiques</i>	49
<i>Analyses démographiques</i>	49
4.3. Résultats	54
<i>Identification des génotypes semblables</i>	54
<i>Fermeture démographique des populations</i>	55
<i>Estimation de l'abondance</i>	55
<i>Estimation de la superficie effective du site d'étude</i>	57
<i>Estimation de la densité de population à partir de l'abondance et de l'aire effective d'étude</i>	57
<i>Estimation de la densité avec la maximisation de la vraisemblance pour des captures-recaptures spatialement explicites (ML SECR) dans le logiciel DENSITY 4.4</i>	57
<i>Estimation de la densité avec la maximisation de la vraisemblance pour des captures-recaptures spatialement explicites (ML SECR) dans le logiciel R</i>	60
<i>Comparaison des estimations de densité obtenues avec les trois approches utilisées : aire effective, ML SECR dans le logiciel DENSITY 4.4 et ML SECR dans le logiciel R</i>	62
<i>Comparaison des estimations d'abondance et de densité obtenues dans cette étude avec celles de Courtois et coll. (2004) pour le site de R2001</i>	63
5. ESTIMATION DE LA DENSITÉ SELON DIFFÉRENTS SCÉNARIOS D'ÉCHANTILLONNAGE PAR SIMULATION	65
5.1. Objectifs et méthodologie des simulations	65
5.2. Résultats des simulations	67
6. DISCUSSION ET RECOMMANDATIONS	69

6.1 Nombre et disposition des stations d'échantillonnage	69
6.2 Efficacité de la méthode de capture et de récolte des poils	70
6.3 Efficacité des appâts et des leurres olfactifs	71
6.4 Période et durée de l'échantillonnage.....	71
6.5 Analyses génétiques	72
6.6 Analyses démographiques	75
<i>Fermeture démographique et géographique des populations d'ours noirs</i>	75
<i>Estimation de la densité de population chez l'ours noir</i>	76
7. CONCLUSION	77
8. BIBLIOGRAPHIE	79

LISTE DES TABLEAUX

Tableau 1. Marqueurs microsatellites dinucléotides utilisés pour identifier les individus dans les populations d'ours noirs. La précision acceptée correspond à la probabilité d'identité (probabilité que deux individus différents aient le même génotype) tenant compte du plus haut niveau d'apparement possible dans la population (PI_{sibs}).....	20
Tableau 2. Description des modèles d'estimation de la taille d'une population fermée selon la prise en compte des effets du temps, de la réponse comportementale à une première capture et de la variation individuelle dans la probabilité de capture.	25
Tableau 3. Modèles d'estimation de la densité pour les populations fermées selon la maximisation de la vraisemblance à partir d'un historique de capture-recapture spatialement explicite. Les modèles incluent un ou plusieurs des effets suivants : le temps (t), la réponse comportementale à une première capture (b) et la variation individuelle dans la probabilité de capture (h).	31
Tableau 4. Nombre et densité (/20 km ²) de stations d'échantillonnage par site d'étude.....	45
Tableau 5. Dates d'installation des stations d'échantillonnage, périodes et durées des sessions d'échantillonnage pour chaque étude.	48
Tableau 6. Marqueurs microsatellites utilisés pour l'identification individuelle pour chaque inventaire.	48
Tableau 7. Paires d'individus semblables, nombre d'individus différents identifiés et nombre de recaptures pour tous les sites d'étude. Les recaptures d'un même individu lors d'une même occasion de capture sont compilées.	54
Tableau 8. Estimation de la taille de population pour R2001 selon les modèles de populations fermées dans MARK. Les estimations et les erreurs types en caractère gras sont celles issues des modèles pour lesquels l'estimation a vraisemblablement échoué.	56
Tableau 9. Estimation de la taille de population pour R2003-E selon les modèles de populations fermées dans MARK. Les estimations et les erreurs types en caractère gras sont celles issues des modèles pour lesquels l'estimation de l'abondance a vraisemblablement échoué.	57
Tableau 10. Superficie effective de l'aire d'étude (km ²) pour tous les sites, selon l'estimateur manuel (bande de 5 056 m), l'ARL/2 et les polygones convexe et concave. Les nombres entre parenthèses correspondent à la largeur (m) de la bande entourant les stations d'échantillonnage.....	59

Tableau 11. Densité d'ours (individu/10 km ²) estimée dans chaque site d'étude avec l'approche $N/\hat{A}$ à partir de l'aire effective d'étude calculée en utilisant la méthode du polygone concave ou du polygone convexe, et proportion des individus capturés une seule fois selon les génotypes.....	59
Tableau 12. Densité d'ours estimée (individu/10 km ²) pour le secteur R2003-E pour chacun des modèles avec DENSITY.....	60
Tableau 13. Densité d'ours estimée (individu/10 km ²) pour le secteur R2001 pour chacun des modèles testés avec l'approche ML SECR dans R.....	61
Tableau 14. Densité d'ours estimée (individu/10 km ²) pour le secteur R2003-E pour chacun des modèles avec l'approche ML SECR dans R.....	62
Tableau 15. Comparaison des estimations de densité calculées à partir des estimations d'abondance et de la superficie de l'aire effective d'étude ($N/\hat{A}$) et celles issues de la maximisation de la vraisemblance de captures-recaptures spatialement explicites (ML SECR) dans DENSITY et dans R. Les erreurs types (S.E.), les intervalles de confiance (IC95 %) et la précision (%) sont aussi présentés pour chaque site d'étude. Les densités portant un astérisque ont été calculées avec une bande entourant les stations de 5 046 m de largeur (établie manuellement).....	64
Tableau 16. Tableau tiré de Courtois et coll. (2004) présentant l'estimation de l'abondance, les erreurs types (S.E.) et les intervalles de confiance (IC95 %) pour les historiques de capture surestimant le nombre probable de recaptures, estimant le plus correctement possible le nombre probable de recaptures et sous-estimant le nombre probable de recaptures (pour plus de détails sur les différents critères de sélection des événements probables de recapture, voir Courtois et coll. [2004]).....	64
Tableau 17. Nombre de captures, de recaptures et de stations conservées pour chacune des bases de données utilisées dans les simulations permettant d'estimer la densité d'ours à partir d'un sous-échantillon des stations.	66
Tableau 18. Estimation de la densité, erreur type et identification du meilleur modèle selon la proportion de stations conservées dans la base de données de R2001.....	67
Tableau 19. Estimations d'abondance obtenues en fonction de différents scénarios de sous-échantillonnage. Tiré de Courtois et coll. (2004).....	75

LISTE DES FIGURES

Figure 1. Exemple d'un piège à poils triangulaire utilisé pour la capture de poils d'ours noirs. L'appât est placé dans une chaudière trouée fixée à un arbre. (Crédit photo : MFFP)	13
Figure 2. Représentation de la moitié de la distance asymptotique maximale entre un évènement de capture et un évènement de recapture	28
Figure 3. Représentation d'une aire effective d'étude utilisant a) un polygone concave et b) un polygone convexe.....	52
Figure 4. Relation entre la précision de l'estimation de densité (étendue de l'intervalle de confiance divisée par l'estimation) et le pourcentage (%) de stations conservées pour l'estimation de la densité pour R2001.....	68

LISTE DES ANNEXES

Annexe 1. Modalité d'échantillonnage et résultats partiels d'études de CMR utilisant la reconnaissance individuelle par génotypage des poils réalisées chez l'ours noir.	93
Annexe 2. Cartes des secteurs d'inventaire de l'ours noir en Abitibi-Témiscamingue de 2001 à 2003.....	95
Annexe 3. Mode d'emploi pour le logiciel IDENTITY 4.	101
Annexe 4. Mode d'emploi pour le logiciel CloseTest.....	104
Annexe 5. Mode d'emploi pour MARK.....	108
Annexe 6. Mode d'emploi pour DENSITY 4.4.	116
Annexe 7. Mode d'emploi pour le paquet SECR dans le logiciel R 2.12.2.	120

1. INTRODUCTION

L'ours noir (*Ursus americanus*) a longtemps été perçu comme une espèce nuisible au Québec, s'attaquant au bétail et saccageant les récoltes. Avant les années 1980, son exploitation se résumait au contrôle de la déprédation, au piégeage pour sa fourrure et aux récoltes opportunistes des chasseurs de cerfs et d'orignaux (Lamontagne et coll., 2006). Bien qu'un permis légal de chasse ait été établi en 1976, la gestion de l'exploitation de l'ours noir ne suscitait alors que très peu d'intérêt (Gignac et coll., 2008). Ce n'est qu'au début des années 1980 que les gestionnaires de la faune ont débuté la mise en valeur de l'ours noir. On lui a alors attribué le statut de gros gibier. Face à une croissance importante de la vente de permis de chasse et des prises enregistrées, on a décidé de s'attarder davantage à la gestion des populations d'ours noirs au Québec. Ainsi, depuis 1984, la récolte annuelle d'ours est suivie et des plans de gestion sont produits (Lamontagne et coll., 2006). Grâce aux statistiques de récolte annuelles, il a été possible d'établir que le nombre d'ours noirs récoltés par la chasse et le piégeage au Québec a triplé de 1984 à 1995 (Jolicoeur et Lamontagne, 1997). Après une baisse attribuable à l'interdiction de la vente de la vésicule biliaire et à l'abandon de la chasse automnale en 1998, la récolte d'ours noirs s'est stabilisée à environ 4 500 ours au milieu des années 2000 (Lamontagne et coll., 2006).

L'ours noir représente sans aucun doute une ressource d'une grande valeur économique pour le Québec. Cependant, au cours des années 1990, des indices suggéraient que certaines populations étaient exploitées à leur niveau maximal et même surexploitées. En effet, entre 1984 et 1995, on a assisté à une augmentation de la proportion de femelles et de l'âge moyen des mâles et des femelles dans la récolte, ainsi qu'à une baisse de l'indice de productivité des femelles. En parallèle, la perception de l'ours noir comme espèce nuisible ne s'est guère améliorée suite à l'augmentation des signalements d'ours importuns et à de rares cas d'attaque de l'ours envers l'homme (Lamontagne et coll., 2006). Ces apparentes contradictions mettent en lumière le manque de connaissances sur les populations d'ours noirs du Québec. Suite à l'analyse des niveaux d'exploitation de l'espèce au Québec, le ministère des Ressources naturelles et de la Faune (MRNF) a souligné l'importance de maintenir les populations à des niveaux biologiquement et socialement acceptables, en plus

d'optimiser les retombées économiques liées à son exploitation (Lamontagne et coll., 2006). Pour ce faire, les quotas de prélèvement doivent représenter fidèlement le nombre maximal d'individus qu'il est possible de récolter sans mettre en péril la population à court et à long terme. Il est donc primordial de pouvoir estimer les densités et la taille des populations de façon précise et régulière, afin d'être en mesure de réagir rapidement et efficacement à des fluctuations de population.

L'approche qui est actuellement la plus répandue pour estimer la taille et la densité des populations d'ursidés est la méthode de capture-marquage-recapture (CMR). Dans ce type d'étude, la proportion d'individus capturés et marqués, puis des individus recapturés lors des sessions d'échantillonnage suivantes, permet d'estimer le nombre d'individus présents dans la population (Seber, 2002). Il est toutefois difficile de capturer les ursidés, car ils vivent à de faibles densités et ils ont des comportements solitaires et élusifs (Boulanger et coll., 2002; Bellemain et coll., 2005; Mainguy et Bernatchez, 2007). De plus, les captures directes peuvent induire un stress non négligeable à l'animal (Cattet et coll., 2003) en plus de présenter des risques pour la sécurité des techniciens et la survie de l'animal lui-même (Amstrup et Beecham, 1976; Lemieux, 1985). Plusieurs variantes de la méthode traditionnelle de CMR ont été développées. Parmi ces méthodes, on retrouve la réobservation des individus avec des caméras sensibles aux mouvements (Grogan et Lindzey, 1999). Il s'avère cependant très difficile, voire impossible de différencier les individus les uns des autres (Foster et Harmsen, 2012). Il existe également d'autres méthodes qui consistent à localiser par voie terrestre ou aérienne les individus portant des colliers émetteurs (Solberg et coll., 2006), à compiler les événements de recapture des ours récoltés par la chasse (Lemieux et Desrosiers, 2002) ou à analyser le matériel biologique récolté suite au marquage chimique des individus (Jolicoeur et Lemieux, 1984; Jolicoeur, 2004; Peacock et coll., 2011). Ainsi, plusieurs des méthodes d'estimation de la densité requièrent la manipulation initiale des individus, ce qui maintient les risques et les coûts associés aux captures directes. De plus, les méthodes de localisation aérienne sont souvent inadéquates puisqu'elles sont coûteuses et qu'il est difficile d'observer des ours en milieu forestier (Solberg et coll., 2006).

Au Québec, Jolicoeur et Lemieux (1984) ont testé une méthode d'évaluation de la densité à l'aide d'un traceur radioactif. Cette méthode consistait à capturer des individus à l'aide de pièges à pattes, à marquer ces individus avec une étiquette et à leur injecter un marqueur radioactif, soit le $^{65}\text{Zinc}$ ou le $^{134}\text{Césium}$. Les fèces d'ours dans l'aire d'étude étaient par la suite récoltées et la densité d'ours était estimée à partir de la proportion de fèces radioactives récoltées. Avec cette méthode, Jolicoeur (2004) a pu obtenir des taux de recapture (fèces radioactives récoltées) de l'ordre de 15 à 59 %, ce qui a permis d'estimer que la densité d'ours noirs dans les différentes régions étudiées variait de 0,10 à 0,55 ours/km². Jugée satisfaisante par Jolicoeur et Lemieux (1990), cette méthode a été abandonnée puisqu'elle impliquait beaucoup de démarches administratives en plus de susciter la crainte chez la population. Jolicoeur (1996) a aussi décrit une autre méthode alternative aux CMR traditionnelles : le marquage à la tétracycline. Cette méthode se base sur le marquage des individus à l'aide de la tétracycline, un antibiotique qui fait jaunir les os et les dents. La tétracycline est cachée dans des appâts distribués sur le territoire échantillonné. La taille et la densité de la population sont ensuite estimées à partir du nombre d'appâts consommés par un ours noir après 10 à 14 jours (événements de capture) et du nombre d'ours récoltés par la chasse ou le piégeage dont les os ou les dents étaient jaunies (événements de recapture). Cette méthode a été jugée inintéressante puisque la précision des estimations d'abondance et de densité était fortement dépendante de la pression de chasse et de piégeage, et de la coopération des chasseurs et des piégeurs.

À la fin des années 1990, une nouvelle approche d'estimation des populations d'ours a été développée et testée (p. ex., Taberlet et coll., 1997; Woods et coll., 1999; Boulanger et coll., 2002; Romain-Bondi et coll., 2004; Roy, Albert et Bernatchez, 2007). Cette méthode est non invasive, c'est-à-dire qu'elle ne nécessite pas la manipulation d'animaux vivants, et elle se base sur la récolte de matériel biologique (fèces ou poils) lors de différentes sessions d'échantillonnage pour constituer des événements de capture et de recapture. L'identification des individus est possible grâce à l'ADN contenu dans le matériel récolté. Tout comme la méthode traditionnelle de CMR, elle permet d'estimer la taille et la densité des populations, mais elle ne nécessite pas la capture directe des individus. Dans une optique de gestion efficiente de l'espèce, il est essentiel de sélectionner une méthode

d'échantillonnage des populations qui tient compte de l'écologie de l'espèce visée et qui offre un bon compromis entre la précision des estimations, les coûts de réalisation et le temps nécessaire pour l'échantillonnage (De Barba et coll., 2010). Il est donc d'un grand intérêt de s'attarder au potentiel de la méthode de CMR par identification individuelle à l'aide d'analyses génétiques pour un suivi à moyen ou long terme des populations. Bien que plusieurs méthodes de CMR non invasives existent, nous nous attarderons seulement à la méthode de CMR utilisant les poils comme matériel génétique, puisque c'est celle qui offre le plus d'avantages pour l'étude des populations d'ours.

2. OBJECTIFS

Les objectifs de ce rapport sont les suivants :

- 1) Faire une revue de la littérature sur la méthode de CMR permettant l'identification individuelle par analyses génétiques;
- 2) Analyser et interpréter les résultats des études expérimentales réalisées avec cette méthode dans la région de l'Abitibi-Témiscamingue;
- 3) Formuler des recommandations quant à l'utilisation de cette approche pour estimer l'abondance et la densité des populations d'ours noirs;
- 4) Proposer un protocole d'inventaire des populations d'ours noirs du Québec.

Dans la section 3 du présent rapport, nous décrivons les techniques d'inventaire par l'approche de capture-marquage-recapture, plus précisément celle de reconnaissance individuelle par analyse génétique des poils (design et disposition des stations d'échantillonnage, analyses génétiques, techniques de modélisation statistiques, etc.). La section 4 fournit la description, l'analyse et les résultats des études expérimentales d'inventaire de l'ours noir réalisées entre 2001 et 2003 en Abitibi-Témiscamingue. La section 5 présente les résultats de simulations visant à estimer les densités à partir des données existantes en utilisant différents scénarios d'échantillonnage, alors que dans la section 6, nous discutons des résultats des études expérimentales et formulons des recommandations basées sur la revue de littérature et les résultats des projets réalisés en Abitibi. Finalement, nous proposons un protocole d'échantillonnage pour l'estimation des

densités de l'ours noir au Québec à la section 7. Ce projet a pu être réalisé grâce à un financement provenant du Réinvestissement dans le domaine de la faune (RDF).

3. REVUE DE LITTÉRATURE

3.1. Techniques d'inventaire de l'ours noir utilisant une approche de CMR

Description de la méthode de CMR avec identification individuelle par génotypage

Une approche aujourd'hui bien répandue pour estimer les populations d'ours noirs est la méthode de CMR avec identification individuelle par génotypage (Lemieux et Desrosiers 2002; Roy et coll., 2007). Cette méthode s'est avérée très utile pour estimer les populations d'espèces rares ou élusives comme les ursidés dans plusieurs régions du monde. Elle permet de différencier les individus provenant d'une même population par l'analyse génétique de matériel biologique récolté sur le terrain (Waits et Paetkau, 2005). Cette approche est particulièrement intéressante, car elle ne nécessite pas la manipulation des animaux. La popularité croissante de cette méthode est, entre autres, reliée aux développements récents des techniques d'analyse génétique avec les chaînes en réaction par polymérase [PCR] (Higuchi et coll., 1988; Saiki et coll., 1988). Grâce aux PCR, il est maintenant possible d'obtenir suffisamment d'ADN à partir de fèces ou de quelques poils pour différencier les individus capturés et reconnaître les individus qui ont déjà été capturés lors des sessions d'échantillonnage précédentes. Dans une étude de CMR par génotypage, la récolte de matériel biologique lors de sessions d'échantillonnage successives (p. ex., des poils récoltés à l'aide de fils barbelés) constitue les événements de capture et de recapture. Le marquage est permanent puisqu'il est basé sur le bagage génétique des individus. Cette méthode a fait ses preuves avec diverses espèces telles que l'ours noir (p. ex., McCall, 2002; Boersen et coll., 2003; Triant et coll., 2004; Roy et coll., 2007; Tredick et coll., 2007; Immell et Anthony, 2008; Settlage et coll., 2008; Robinson et coll., 2009; Gardner et coll., 2010), le grizzly [*Ursus arctos horribilis*] (Taberlet et coll., 1997; Poole et coll., 2001; Mowat et coll., 2004; Romain-Bondi et coll., 2004; Kasworm et coll., 2007), le loup gris [*Canis lupus*] (Creel et coll., 2003), la martre [*Martes americana*] (Foran et coll., 1997), le blaireau [*Meles meles*] (Balestrieri et coll., 2010), la baleine à bosse [*Megaptera*

novaeangliae] (Palsboll et Allen, 1997) et la tortue peinte [*Chrysemys picta*] (Pearse et coll., 2001).

Tout comme les études de CMR traditionnelles, les techniques de CMR par génotypage supposent que :

- 1) la population est fermée démographiquement et géographiquement, c'est-à-dire qu'il n'y a aucune naissance, mortalité, entrée ou sortie d'individus de l'aire d'étude durant la période d'échantillonnage;
- 2) tous les individus ont la même probabilité de capture et cette probabilité reste constante durant la période d'échantillonnage;
- 3) le marquage doit être permanent et doit permettre la reconnaissance des individus sans erreur;
- 4) le marquage ne doit pas influencer le sort de l'animal, en particulier ses chances de recapture.

Ces suppositions sont rarement respectées en totalité. Les conséquences de la violation d'une de ces suppositions sur les estimations d'abondance et de densité sont discutées plus loin. Il existe des techniques de modélisation pour estimer l'abondance qui permettent de diminuer l'importance d'une violation de ces suppositions et des éléments à considérer lors de l'établissement des protocoles expérimentaux afin de minimiser le risque de violation.

La méthode de CMR par génotypage présente de nombreux avantages par rapport aux autres méthodes traditionnelles d'estimation des populations. D'une part, elle ne nécessite pas la manipulation des animaux, ce qui permet de diminuer la fréquence des visites sur le terrain et, par conséquent, les coûts d'un projet (Solberg et coll., 2006). D'autre part, elle est considérée comme plus éthique, car elle est plus sécuritaire pour les animaux et les humains (Taberlet et Luikart 1999). Cette technique permet aussi d'obtenir un grand

nombre d'évènements de capture et de recapture, ce qui est essentiel à l'estimation des populations et est généralement difficile à obtenir pour des espèces rares et élusives comme les ursidés (Mills et coll., 2000; Miller et coll., 2005). Le fait d'utiliser le bagage génétique des individus permet d'avoir un marquage permanent, ce qui diminue les risques de biais pour l'estimation des populations (p. ex., McDonald et coll., 2003) et représente un avantage considérable pour des suivis à long terme. Solberg et coll. (2006) ont comparé trois méthodes d'estimation de population : l'observation directe des femelles avec des ours, une étude de CMR de femelles en œstrus munies de collier émetteur par inventaire aérien et une étude de CMR par génotypage à partir de fèces récoltées de façon opportuniste. Les auteurs ont démontré que la technique de CMR avec génotypage était la plus efficace, la plus fiable et la moins dispendieuse pour estimer la taille d'une population d'ours bruns. Cependant, il existe également des avantages à utiliser des poils plutôt que des fèces comme matériel biologique. En effet, on retrouve beaucoup moins d'ADN dans les fèces que dans la racine des poils, et puisque l'ADN se dégrade très rapidement, cela augmente le risque d'échec ou d'erreur dans l'identification des individus (Panasci et coll., 2011). Aussi, il est beaucoup plus facile de récolter des poils sur le terrain à l'aide de pièges à poils et d'un leurre que de récolter des fèces de façon opportuniste ou systématique.

Par ailleurs, cette méthode s'est avérée pertinente pour une multitude d'applications en écologie (voir Waits et Paetkau, 2005, pour une liste exhaustive des applications chez plusieurs espèces animales). Par exemple, ce type d'échantillonnage peut servir à la détection d'espèces rares (p. ex., Taberlet et coll., 1997), à la détermination de la variabilité génétique au sein d'une population (p. ex., Boersen et coll., 2003; Triant et coll., 2004), à l'isolement géographique de populations (p. ex., Robinson et coll., 2007; Brown et coll., 2009; Harris et coll., 2010) ou aux déplacements des individus (ex. Roy et coll., 2007). L'identification génétique des individus peut aussi avoir des applications légales, telles que la détermination de la provenance d'organes vendus au marché noir (Waits et coll., 1999).

Caractéristiques de l'aire d'étude et disposition des blocs pour un inventaire de CMR par géotypage

Il est souvent difficile, voire impossible, de couvrir entièrement le territoire pour lequel on veut estimer la taille d'une population lors d'un échantillonnage par CMR, à cause de contraintes de temps ou d'argent. Or, extrapoler les estimations de population et de densité obtenues dans de petites aires études à une région plus grande peut être discutable (Boulanger et coll., 2004). Il faut donc choisir une aire d'étude suffisamment petite pour être totalement couverte par le dispositif expérimental ou sélectionner quelques blocs d'étude représentatifs de la diversité du territoire (p. ex., Roy et coll., 2007). Le choix du site d'étude doit minimiser les chances que des individus y entrent ou en sortent durant la période d'échantillonnage, ce qui pourrait biaiser les résultats (fermeture géographique de la population). Il est possible de délimiter la grille d'échantillonnage en fonction de barrières géographiques naturelles telles que des autoroutes, des villes ou des chaînes de montagnes. Ces barrières physiques permettent aussi de limiter la probabilité que les domaines vitaux des individus chevauchent la limite de la grille (effet de bordure). Dans ce cas, les individus peuvent seulement être capturés lorsqu'ils se trouvent dans la section de leur domaine vital située dans l'aire d'étude. Toutefois, il est souvent difficile de délimiter l'aire d'étude en fonction de barrières géographiques. Il est cependant possible de réduire l'effet de bordure, en augmentant le ratio entre la superficie de la grille d'échantillonnage et la taille moyenne des domaines vitaux (Boulanger et coll., 2002). L'utilisation d'une aire d'étude effective qui inclut la portion des domaines vitaux se trouvant à l'extérieur des limites de l'aire d'étude permet de minimiser l'effet de bordure. Il existe plusieurs estimateurs permettant de calculer l'aire effective (voir la section 4.2). Finalement, la superficie de la grille d'échantillonnage sera également déterminée par de nombreux facteurs, tels que le nombre de stations nécessaires pour obtenir l'intensité d'échantillonnage voulue, leur disposition et leur accessibilité.

Nombre et disposition des stations d'échantillonnage

Le nombre et la disposition des stations d'échantillonnage sont des paramètres déterminants pour le succès d'une étude de CMR. En effet, une des suppositions de base des études de CMR est la probabilité égale de capture pour tous les individus. Il faut donc que le design expérimental permette que tous les individus de la population aient accès à au moins une station d'échantillonnage. Le nombre optimal de pièges à installer dépend de la superficie moyenne des domaines vitaux du sexe le moins mobile [les femelles dans le cas des ursidés] (Rogers, 1987; Samson, 1996; Jolicoeur, 1997a; Alt et coll., 1980; Samson et Huot, 1994; Roy et coll., 2007) et de la capacité des équipes de terrain en fonction des restrictions budgétaires. La superficie des domaines vitaux est fortement dépendante de la qualité de l'habitat, plus précisément de la quantité et de la diversité de la nourriture qu'on y retrouve (Rogers, 1987; Samson et Huot, 1994). Chez l'ours noir, la superficie des domaines vitaux des mâles peut être cinq fois plus élevée que celle des femelles; elle peut atteindre en moyenne 490 km² pour les mâles (Massé, 2012) et 95 km² pour les femelles (Rogers, 1987; Samson et Huot, 1994; Jolicoeur, 1997a). La superficie moyenne des domaines vitaux répertoriée parmi différentes études en Amérique du Nord est toutefois inférieure à ces valeurs, variant de 8 à 19 km² pour les femelles et de 21 à 166 km² pour les mâles (Samson, 1996). Cependant, les ours noirs semblent aussi concentrer leurs activités dans certaines zones de leur domaine vital, appelées zones d'utilisation intensive (Rogers, 1987; Pacas et Paquet, 1994; Leblanc et Huot, 2000). Une étude réalisée dans le parc national Forillon a démontré que les femelles utilisaient jusqu'à cinq zones d'utilisation intensives, chacune d'elles couvrant de 1,8 à 2,9 km². Une autre étude dans le parc national de la Mauricie a démontré des zones d'utilisation intensive de 10 à 16 km² composées de milieux perturbés et de forêts de feuillus tolérants, deux milieux propices à l'alimentation pour l'ours noir (Samson et Huot, 1994).

Une étude sur l'ours noir réalisée par Roy et coll. (2007) a permis de mettre en évidence l'importance de l'effort d'échantillonnage, ici représenté par le nombre de stations, dans l'estimation de la taille et de la densité d'une population. Pour six blocs couvrant chacun

500 km², ils ont placé une station par parcelle de 25 km². Puisque la superficie des domaines vitaux des femelles était évaluée à 45 km² dans la région de l'étude, des parcelles de 25 km² auraient dû permettre à toutes les femelles d'avoir accès à au moins une station. En effet, avec un tel scénario, il devrait y avoir en moyenne 1,8 station dans le domaine vital de chaque femelle. Toutefois, deux parcelles par blocs ont été munies de cinq stations d'échantillonnage. Les auteurs ont montré que plusieurs nouvelles femelles ont été identifiées seulement aux quatre stations supplémentaires, ce qui suggère que toutes les femelles n'avaient pas accès à une station d'échantillonnage dans les parcelles contenant une seule station. Ainsi, ils ont conclu que l'effort d'échantillonnage était insuffisant malgré la sélection de parcelles plus petites que la superficie des domaines vitaux, suggérant que cette dernière avait été surestimée. Settlage et coll. (2008) ont obtenu des taux de capture et de recapture considérés comme suffisants pour estimer la taille des populations en utilisant des parcelles nettement plus petites que la superficie moyenne des domaines vitaux des femelles. Dans le sud des Appalaches, aux États-Unis, ils ont établi une station d'échantillonnage par parcelle de 2,5 km² et une station par parcelle de 5,8 km² dans des secteurs où la superficie des domaines vitaux était évaluée à 8,4 km² et à 14,0 km² respectivement.

Pour s'assurer que tous les individus ont la même accessibilité aux stations d'échantillonnage, Otis et coll. (1978) ont suggéré que quatre stations par surface correspondant à la taille du domaine vital moyen étaient nécessaires. La densité moyenne des stations dans les études de CMR par génotypage chez l'ours noir était de 1 station/7,7 km² dans la littérature consultée (annexe 1), ce qui correspondrait à quatre stations en moyenne par domaine vital dans une région où les domaines vitaux des femelles seraient d'environ 30 km². Une autre approche pour maximiser la probabilité de capture des individus est de déplacer les stations d'échantillonnage à l'intérieur des parcelles entre les sessions (White et coll., 1982; voir l'annexe 1 pour des exemples), ce qui peut favoriser l'accès à une station aux femelles dont les territoires ne chevauchent pas l'emplacement initial de la station, mais peut complexifier considérablement le plan d'échantillonnage et la logistique des travaux sur le terrain.

La disposition des stations d'échantillonnage peut être systématique ou non à l'intérieur de l'aire ou des blocs d'étude. Dans le cas d'une disposition non systématique des stations, la localisation de celles-ci peut être choisie subjectivement en fonction de la qualité ou de l'accessibilité de l'habitat, tout en respectant un espacement minimum prédéterminé des stations (Boersen et coll., 2003). Dans ce cas, il faut toutefois être prudent quant à l'interprétation des résultats d'abondance et de densité, puisque ces derniers risquent d'être surestimés par rapport à ceux d'habitats moins propices aux ours. La disposition des stations d'échantillonnage peut aussi être aléatoire, ce qui peut être réalisé aisément à l'aide de systèmes d'information géographique [SIG] (p. ex., ArcGIS) en respectant certains critères comme l'accessibilité et la répartition équitable des stations dans différents types d'habitats (Settlage et coll., 2008).

Pièges à poils, appâts et leurres olfactifs aux stations d'échantillonnage

Chez les ursidés, la collecte de poils à l'aide de pièges à poils entourant un site appâté ou avec un leurre olfactif s'est avérée un moyen très efficace d'obtenir du matériel génétique (Taberlet et Bouvet, 1992). Les animaux se frottent au piège constitué de fils barbelés qui retiennent alors quelques poils qui possèdent encore leur follicule, offrant ainsi une plus grande quantité d'ADN pour les analyses génétiques que les poils seuls (Higuchi et coll., 1988; Gagneux et coll., 1997). Quelques types de pièges à poils ont été testés sur le terrain, comme le piège de barbelés triangulaire et l'enclos de fils barbelés. Le piège de fils barbelés triangulaire (figure 2) est le plus rapide d'installation et le moins dispendieux, puisqu'il est plus simple et requiert moins de matériel (Woods et coll., 1999; Lemieux et Desrosiers 2002; Roy et coll., 2007). Cependant, un problème important avec le piège triangulaire est qu'il n'offre qu'une seule voie d'accès à l'appât pour les ours, ce qui augmente le risque de contamination des échantillons. En effet, deux ours qui visitent la même station laisseront des poils sur les mêmes pointes du fil barbelé, et leur identification génétique individuelle sera impossible (Courtois et coll., 2004). Immell et Anthony (2008) ont conçu un autre type de piège à poils où l'ours déclenche le piège lors de sa montée dans un arbre pour atteindre l'appât. Ainsi, un seul échantillon de poils est récolté par piège par session d'échantillonnage, évitant le regroupement de poils provenant de deux ou plusieurs

ours dans un seul échantillon. Cependant, cette méthode nécessite un nombre élevé de pièges pour maximiser le nombre de captures et de recaptures, puisqu'une seule capture (ou recapture) est possible par piège par session d'échantillonnage (Immell et Anthony 2008).

Plusieurs types d'appâts peuvent être utilisés pour inciter les ours à passer sous les fils barbelés. Par exemple, Lemieux et Desrosiers (2002) ont utilisé des gâteries composées de retailles de gâteaux commerciaux, de miel, de confiture et de mélasse. D'autres études ont utilisé des sous-produits du phoque (Roy et coll., 2007) ou un mélange de retailles de viande, de poisson, de pâtisseries et de sirop (Boersen et coll., 2003). Les gâteries ne nécessitent pas de réfrigération lors des travaux de terrain et peuvent se conserver plus longtemps que des morceaux de viande. On évite ainsi la prolifération des vers aux installations. Elles sont toutefois peu odorantes et donc peu attractives pour les ours, de sorte qu'il faut les combiner avec un leurre olfactif. Selon Woods et coll. (1999), il est cependant préférable de ne pas fournir des appâts alimentaires pour éviter que les ours s'habituent à cette source de nourriture et qu'ils développent un comportement « trappophile ». Il existe plusieurs types de leurres olfactifs pour attirer les ours à partir de grandes distances : du sang (vache, bœuf, phoque), de la chaire de phoque en décomposition dans des chaudières et dont la partie liquide est aussi vaporisée sur les arbres autour du piège (Roy et coll., 2007), du vaporisateur commercial aromatisé à la fraise et à la framboise, mélangé à du propylène glycol et de l'huile essentielle d'anis (Lemieux et Desrosiers 2002) ou un vaporisateur à base de fertilisant ou d'huile de poisson (Woods et coll., 1999; Poole et coll., 2001; Romain-Bondi et coll., 2004). Le leurre olfactif doit pouvoir attirer les ours présents dans les parcelles d'échantillonnage et rester odorant entre les visites aux pièges malgré les intempéries.



Figure 1. Exemple d'un piège à poils triangulaire utilisé pour la capture de poils d'ours noirs. L'appât est placé dans une chaudière trouée fixée à un arbre. (Crédit photo : MFFP).

Période et durée de l'échantillonnage

La période d'échantillonnage, la fréquence des visites aux stations et la durée des occasions de capture sont aussi des paramètres importants pour la réussite d'une étude de CMR par génotypage des poils. L'échantillonnage doit être effectué à une période de l'année où tous les ours sont actifs afin de maximiser et d'uniformiser les chances de capture entre les individus. Rogers (1987) a suggéré que les sessions d'échantillonnage pour l'ours noir devaient s'effectuer en juin, alors que les oursons de l'année en mauvaise condition sont déjà morts et que les longues excursions pour la recherche de nourriture n'ont pas encore débuté. Gignac et coll. (2008) ont suggéré de procéder aux échantillonnages pour l'ours noir alors que les petits fruits ne sont pas encore apparus, soit entre le début juin et la fin juillet, afin de profiter de la rareté de la nourriture en nature pour augmenter le taux de visite aux stations d'échantillonnage. Un des postulats de base des études de CMR est que la population soit fermée durant l'échantillonnage, c'est-à-dire qu'il n'y ait pas de

naissance, de mortalité, d'émigration ou d'immigration dans l'aire d'étude. Il est donc important que l'échantillonnage soit effectué à l'extérieur des saisons de chasse et de piégeage (Roy et coll., 2007).

Au Québec, l'ours noir est chassé et piégé au printemps et à l'automne et l'exploitation par l'homme est la principale cause de mortalité, bien qu'il y ait des mortalités naturelles comme le cannibalisme (p. ex., Samson et Huot, 1994). Le printemps est la saison consacrée pour chasser l'ours noir au Québec; elle débute le 15 mai pour finir généralement le 30 juin. Bien que le piégeage à l'ours noir soit également permis le printemps et l'automne, cette activité se concentre essentiellement le printemps en raison de la qualité de la fourrure.

La durée de l'échantillonnage, quant à elle, doit être assez longue pour obtenir un taux de recapture suffisamment élevé sans être trop longue pour éviter de violer le postulat de fermeture démographique. La majorité des études recensées à l'annexe 1 ont prélevé des échantillons pendant un à deux mois, avec de deux à huit visites de chaque station. L'intervalle de temps entre deux visites d'une même station (occasion de capture) doit être suffisant pour maximiser les chances de capture sans toutefois être trop long, ce qui pourrait compromettre la qualité des échantillons de poils recueillis (visite d'une station par plusieurs ours et dégradation de l'ADN des poils). La plupart des études recensées avaient des occasions de capture espacées de 7 à 21 jours.

En plus du nombre de stations d'échantillonnage, du nombre de sessions et de la durée de l'étude, l'effort d'échantillonnage peut aussi être déterminé en fonction du nombre d'échantillons nécessaires à la viabilité de l'étude. Pour de grandes populations, il est recommandé d'analyser plusieurs échantillons de poils pour favoriser la détection d'évènements de recapture. Or, plus il y a d'échantillons à analyser, plus le coût des analyses génétiques augmente (Waits et Leberg, 2000). Il est avantageux de visiter fréquemment les stations pour maximiser la qualité des échantillons de poils, car l'ADN a tendance à se détériorer rapidement à la chaleur, à la lumière et à l'humidité (Lindahl, 1993). Dans ces conditions, il peut être intéressant d'opter pour le sous-

échantillonnage des échantillons de poils. Cette approche ne diminue pas la fréquence des visites et donc les coûts du travail sur le terrain, mais elle permet une couverture uniforme de l'aire d'étude, à la fois dans le temps et dans l'espace, et diminue les coûts reliés aux analyses génétiques (Boersen et coll., 2003; Settlage et coll., 2008). Cependant, le nombre de captures et de recaptures conservées dans le sous-échantillon doit être suffisamment élevé pour permettre une estimation fiable de la population, ce qui est impossible à déterminer avant l'analyse génétique des échantillons.

Récolte et conservation des échantillons de poils

La récolte des poils se fait à même les fils barbelés aux stations. L'utilisation d'une feuille blanche, ou d'un autre support de couleur pâle que l'on déplace derrière le fil, facilite la détection des poils sur les barbelés (Roy et coll., 2007). Il est possible que plusieurs ours aient visité la station d'échantillonnage durant une même session et que plusieurs échantillons soient présents à la même station. Si on regroupe ces échantillons de poils sans distinction, ceux-ci peuvent révéler un « faux » génotype provenant de poils de plusieurs individus et on ne pourra identifier tous les individus qui ont réellement visité la station (Gagneux et coll., 1997). Il peut en résulter une sous-estimation de la population puisque la station ne permet d'identifier qu'un seul « faux » individu par session alors qu'il y en a plusieurs (Roy et coll., 2007), ou encore une surestimation de la population puisque les « faux » individus identifiés à partir de plusieurs échantillons de poils ne pourront pas être recapturés, augmentant ainsi le nombre d'individus capturés. Une méthode de différenciation des échantillons de poils à la même station a été proposée par Roy et coll. (2007) dans une étude où au moins 22,7 % des stations ont été visitées par plus d'un ours. Tous les poils récoltés d'un même côté du piège de barbelés ont été regroupés dans un même échantillon. Tredick et coll. (2007) ont quant à eux vérifié l'association génétique des poils retrouvés sur deux pointes de barbelés adjacentes et ceux retrouvés à deux pointes d'intervalle. Ils ont déterminé qu'entre 9 et 22 % des échantillons de poils récoltés sur des pointes adjacentes provenaient d'ours différents, alors que cette proportion passait à entre 23 et 55 % pour des échantillons récoltés à deux pointes d'intervalle. Les auteurs suggèrent que ces résultats sont typiques pour des densités élevées d'ours, mais qu'il faut être prudent

dans le regroupement ou l'abandon d'échantillons de poils sur des pointes de barbelés adjacentes pour des populations d'ours à plus faible densité.

Sachant que les poils et leurs follicules contiennent peu d'ADN (Taberlet et coll., 1996) et que ce dernier peut se dénaturer rapidement, il est important d'entreposer les poils récoltés de façon à conserver l'ADN le plus intact possible jusqu'aux analyses génétiques. Aussi, il est facile de contaminer l'ADN recueilli lors de la récolte et de la manipulation. D'ailleurs, presque toutes les études utilisant les poils pour la différenciation des individus mentionnent l'utilisation de gants pour récolter les poils et la stérilisation par le feu des fils barbelés entre chaque session d'échantillonnage pour éviter la contamination du matériel génétique (McCall 2002; Triant et coll., 2004; Roy et coll., 2007). L'entreposage des poils peut se faire de différentes façons. Les poils peuvent être entreposés au sec à la température de la pièce. Ils doivent préalablement être séchés dans des sacs de papier, puis placés dans des sacs hermétiques (Triant et coll., 2004; Roy et coll., 2007). Roon et coll. (2003) ont montré que la combinaison de la dessiccation dans des enveloppes et de la congélation dans des sacs hermétiques à -20°C sans cycle de gel et dégel offrait le meilleur succès d'amplification de l'ADN et d'identification individuelle pour l'ours brun (10 % supérieur à celui des poils seulement séchés). Ils ont aussi remarqué que l'ADN provenant des poils séchés et congelés se détériorait rapidement de six mois à un an après leur collecte. Ils soulignent donc l'importance d'effectuer les analyses génétiques rapidement après la collecte des poils pour maximiser les chances de succès. Enfin, on considère que cinq poils de garde avec des follicules intacts ou dix poils de bourre avec des follicules intacts par échantillon sont suffisants pour amplifier avec succès l'ADN d'ours noirs et obtenir suffisamment d'ADN pour le génotypage (Breton et Dufresne, 2003; De Barba et coll., 2010). Utiliser plus de poils signifie obtenir de plus grandes quantités d'ADN pour les analyses génétiques et de résultats pour l'identification (Goossens et coll., 1998).

Analyse génétique des poils

Pour procéder aux analyses génétiques, l'ADN doit d'abord être extrait des poils recueillis. L'ADN nucléaire est généralement utilisé pour les analyses de différenciation individuelle.

Bien qu'il soit moins abondant que l'ADN mitochondrial (Waits et Paetkau, 2005), il est plus variable et donc plus adéquat pour différencier les individus (Higuchi et coll., 1988). L'ADN provenant de poils est extrait principalement par deux méthodes : en appliquant le protocole Chelex (p. ex., Goossens et coll., 1998; Woods et coll., 1999) ou en utilisant des kits d'extraction commerciaux (p. ex., QIAamp™, Qiagen Inc, Ontario, Canada; Poole et coll., 2001). Les deux techniques sont largement utilisées dans l'extraction de l'ADN des poils d'ursidés. Cependant, Poole et coll. (2001) ont observé un taux élevé d'échec d'extraction de l'ADN avec le protocole Chelex avec des échantillons de plus de six poils. Une fois extrait, l'ADN doit être amplifié de façon à obtenir suffisamment d'ADN pour le génotypage. Les chaînes en réaction par polymérase (Saiki et coll., 1988) permettent d'amplifier les séquences d'ADN à partir d'une quantité initiale infime.

Pour différencier les individus, on profite de la variabilité génétique de certains marqueurs microsatellites dans la population. Un marqueur microsatellite est une séquence de deux à quatre nucléotides qui est répétée un nombre variable de fois. Le nombre de répétitions des nucléotides détermine l'identité de l'allèle. Ainsi, pour un microsatellite donné, il peut exister plusieurs allèles possibles qui sont présents selon différentes fréquences dans la population. La combinaison des allèles présents à chaque microsatellite (formant le génotype) permet de différencier les individus (Jarne et Lagoda, 1996). Le nombre de microsatellites utilisés pour l'identification doit être suffisamment élevé pour permettre la différenciation d'individus apparentés, sans être toutefois trop élevé afin de ne pas engendrer des coûts et des erreurs supplémentaires lors des analyses génétiques (voir la section sur les sources d'erreurs possibles lors du génotypage). Plusieurs auteurs ont souligné l'importance de vérifier la variabilité génétique de plusieurs marqueurs microsatellites dans la population avant d'entreprendre une étude (Taberlet et coll., 1999; Waits et Paetkau, 2005). Plus l'hétérozygotie (variabilité génétique) est grande dans la population, plus il est facile de différencier les individus avec peu de microsatellites. Paetkau et Strobeck (1994) ainsi que Breton et Dufresne (2003) ont obtenu des probabilités de fausse identification suffisamment basses pour juger leurs identifications fiables avec seulement quatre marqueurs microsatellites, soit 0,015 et 6×10^{-6} respectivement, alors que 0,01 est un seuil généralement acceptable pour les études de CMR. Cependant, de sept

à dix microsatellites sont généralement nécessaires pour obtenir une forte probabilité d'identifier correctement les individus (Taberlet et coll., 1997; Woods et coll., 1999; Waits et coll., 2001; Paetkau 2003). Chez l'ours noir, les marqueurs microsatellites dinucléotides (répétitions de deux nucléotides) ayant une bonne variabilité sont déjà connus et largement utilisés (tableau 1). Récemment, des marqueurs tétranucléotides ont été découverts chez l'ours noir, marqueurs qui ont un polymorphisme supérieur aux marqueurs dinucléotides en ce qui a trait à la variété d'allèles possibles (Meredith et coll., 2009; Sanderlin et coll., 2009). Les marqueurs tétranucléotides actuellement connus sont A2, A107, B1, B2, B7, B8, B101, B103, B105, B117, B125, B131, C8, C11, C115, C116, D1a, D2, D3, D11, D102, D103, D112, D113, D114, D116, D118, D123 (Meredith et coll., 2009).

La capacité d'identification des individus de la population étudiée peut être quantifiée à l'aide d'un calcul de la probabilité d'identité (PI) qui illustre la probabilité que deux individus différents aient le même patron génotypique pour les marqueurs microsatellites utilisés. Il faut d'abord calculer la probabilité pour chaque locus (équation 1) et ensuite multiplier ces probabilités entre elles pour chaque locus utilisé.

$$PI = \sum p_i^4 + \sum (2p_i p_j)^2 \quad (\text{Équation 1})$$

où p_i et p_j sont les fréquences alléliques pour les allèles i et j . Puisqu'il est plus probable d'obtenir des échantillons de poils d'individus apparentés (mère et oursons ou chevauchement des domaines vitaux entre individus apparentés), la probabilité d'identité calculée à partir des fréquences alléliques de la population est sous-estimée et l'utilité de PI dans ce contexte est moins intéressante (Waits et Paetkau, 2005). Toutefois, un calcul tenant compte de l'apparentement possiblement élevé dans certaines populations, PI_{sibs} (équation 2), a récemment été développé. Celui-ci illustre le plus haut niveau possible d'apparentement dans une population (Taberlet et coll., 1999; Waits et coll., 2001).

$$PI_{sibs} = 0.25 + (0.5 \sum p_i^2) + \left[0.5 (\sum p_i^2)^2 \right] - (0.25 \sum p_i^4) \quad (\text{Équation 2})$$

Les valeurs de PI calculées à l'aide de la formule traditionnelle (équation 1) et celles obtenues avec la formule considérant l'apparentement (équation 2) permettent de déterminer les limites inférieure et supérieure de PI dans une population. Elles permettent

aussi de déterminer le nombre de marqueurs microsatellites à analyser pour atteindre le niveau de confiance voulu. Par exemple, Taberlet et Luikart (1999) ont déterminé que quatre à huit marqueurs étaient nécessaires pour obtenir une probabilité d'identité de 0,01, alors que ce nombre passait à 14 marqueurs pour une probabilité d'identité de 0,0001.

Tableau 1. Marqueurs microsatellites dinucléotides utilisés pour identifier les individus dans les populations d'ours noirs. La précision acceptée correspond à la probabilité d'identité (probabilité que deux individus différents aient le même génotype) tenant compte du plus haut niveau d'apparement possible dans la population (PI_{sibs}).

Source	Marqueurs microsatellites	N ^{bre} d'allèles/		Hétérozygotie moyenne	Probabilité d'identité	Précision acceptée
		de locus	e			
Paetkau et Strobeck, 1994 (Parc national de la Mauricie)	G1A, G1D, G10B, G10L	4-16	0,73	$2,0 \times 10^{-5}$	N/D	N/D
Paetkau et Strobeck, 1994 (Parc national de Banff)	G1A, G1D, G10B, G10L	4-16	0,80	$1,1 \times 10^{-5}$	N/D	N/D
Paetkau et Strobeck, 1994 (Parc national de Terra Nova)	G1A, G1D, G10B, G10L	4-16	0,36	$4,6 \times 10^{-2}$	N/D	N/D
Woods et coll., 1999 (selon Paetkau et coll., 1995)	G1A, G1D, G10B, G10C, G10L, G10X	6-13	0,81	N/D	$PI_{sibs} < 0,05$	
McCall et coll., 2002	G1A, G10B, G10D, G10H, G10J, G10M, G10X, MU59, M-SUT2	N/D	0,74	N/D	N/D	N/D
Boersen et coll., 2003	G1A, G1D, G10B, G10C, G10L, G10M, G10P, G10X (MU10, MU23, MU50, G10I pour analyses complémentaires)	2-6	0,48	$0,454 \text{ à } 0,831$ $PI_{sibs} = 1,52 \times 10^{-2}$	$PI_{sibs} < 0,04$	
Breton et coll., 2003	G1A, G1D, G10B, G10L	9-18	0,83	$5,3 \times 10^{-4}$	$PI_{sibs} < 0,05$	
Triant et coll., 2004	G1A, G1D, G10B, G10C, G10L, G10M, G10P, G10X	5-8	> 0,40	$2,84 \times 10^{-4}$ $PI_{sibs} = 0,0092$	$PI_{sibs} < 0,05$	
Triant et coll., 2004	G1A, G1D, G10B, G10C, G10L, G10M, G10P, G10X	5-8	> 0,55	$5,23 \times 10^{-6}$ $PI_{sibs} = 0,0013$	$PI_{sibs} < 0,05$	
Dreher et coll., 2007	G10D, G10L, G10X, MU50, MU59	7-14	0,75	$2,68 \times 10^{-6}$ $PI_{sibs} = 0,00865$	N/D	
Robinson et coll., 2007	G1A, G1D, G10B, G10C, G10L, G10M, G10P, G10J, G10O, CXX20, MU15, MU23, MU50	N/D	0,67	$6,08 \times 10^{-14}$ $PI_{sibs} = 7,56 \times 10^{-6}$	N/D	
Roy et coll., 2007	G1D, G10H, G10L, G10M, G10P, MU09, MU10, MU15, MU23, MU50	9-20	0,88	$PI_{sibs} = 5,16 \times 10^{-12}$	$PI_{sibs} < 0,01$	

Tableau 1 (suite)

Source	Marqueurs microsatellites	N^{bre} d'allèles/ de locus	Hétérozygotie moyenne	Probabilité d'identité	Précision acceptée
Tredick et coll., 2007 (St. John's)	G1A, G1D, G10B, G10H, MU50, MU59	N/D	N/D	N/D	$PI_{sibs} < 0,01$
Tredick et coll., 2007 (Pocosin Lakes National Wildlife Refuge)	G1A, G1D, G10H, G10J, G10L, MU50	N/D	N/D	N/D	$PI_{sibs} < 0,01$
Immell et coll., 2008	G10B, G10H, G10I, G10M, G10P, MU59	N/D	N/D	N/D	N/D
Settlage et coll., 2008	G1A, G1D, G10B, G10C, G10L, G10M, G10P, G10X, Mu23	N/D	N/D	N/D	N/D
Brown et coll., 2009	G1A, G1D, G10*, G10B, G10H, G10J, G10L, G10P, G10X, MU50, MU59, M- SUT8	3-16	0,53	N/D	N/D

La plupart des études sur les ursidés ont considéré que des valeurs de PI_{sibs} de 0,01 à 0,05 étaient suffisantes pour obtenir un niveau de précision d'attribution de génotype acceptable (Mills et coll., 2000; voir le tableau 1 pour des exemples). Toutefois, Waits et Paetkau (2005) ont souligné l'importance de choisir ce seuil de précision en fonction du nombre de paires d'individus qui sont susceptibles d'être échantillonnés. Dans le cas d'une population isolée, et donc pour laquelle il est très probable de capturer des individus très apparentés, il n'est pas recommandé d'utiliser des calculs basés sur la fréquence de certains allèles dans la population (p. ex., PI_{sibs}) comme estimation de la capacité de bien discriminer les génotypes des individus (Kasworm et coll., 2007; Proctor et coll., 2010).

Sexage des individus

Il est aussi possible, à l'aide du matériel génétique contenu dans les poils récoltés, de déterminer le sexe des individus. Sexer les individus peut permettre d'évaluer le rapport des sexes (p. ex., Taberlet et coll., 1993) et de détecter des différences dans la probabilité de capture entre les sexes (p. ex., Roy et coll., 2007), et offre un critère supplémentaire dans l'identification des individus (Sloane et coll., 2000). Il existe actuellement différentes techniques pour sexer les individus à partir de leurs poils. La plupart d'entre elles se basent sur la détection de la présence d'un gène présent sur le chromosome Y, donc uniquement chez les mâles (Taberlet et coll., 1993; Breton et Dufresne 2003). Toutefois, cette technique est souvent sujette à la contamination puisque les amorces utilisées pour amplifier le gène en question peuvent aussi être présentes chez l'humain. Elle peut aussi échouer à amplifier le gène chez un mâle (Waits et Paetkau 2005). Roy et coll. (2007) ainsi que Poole et coll. (2001) se sont quant à eux basés sur le protocole de Yamamoto et coll. (2002) pour sexer des ours noirs. Ce protocole utilise le gène de la protéine amélogénine présente chez les deux sexes, mais dont la taille allélique diffère entre les sexes. Cette technique s'est avérée efficace malgré les faibles quantités d'ADN utilisées. Il est à noter que les méthodes de sexage actuellement connues sont exposées à d'importants biais si l'échantillon contient les poils de deux individus (ou plus) de sexes opposés (Roy et coll., 2007).

Identification des génotypes identiques

Plusieurs approches sont possibles pour déterminer les génotypes identiques selon la qualité des analyses génétiques. Il est possible que les analyses génétiques ne permettent pas de déterminer le génotype correctement avec tous les loci utilisés. Dans cette situation, on peut déterminer le nombre minimal de loci à utiliser pour réaliser l'appariement. Par exemple, un génotype complet sur trois loci pourrait être considéré comme identique à un génotype complet sur quatre loci, à condition que les trois loci communs aux deux génotypes soient identiques. Bien qu'il soit possible de déterminer d'autres critères de discrimination, tels que la distance maximale entre deux événements de capture ou la distance génétique maximale entre deux allèles semblables pour déterminer les paires d'individus semblables, ces méthodes requièrent un examen visuel des données et elles sont donc plus subjectives (Courtois et coll., 2004). Les logiciels IDENTITY 1.0 (Wagner et Sefc, 1999) et CERVUS 3.0 (Kalinowski et coll., 2007) permettent d'identifier rapidement les paires d'individus identiques, en plus de fournir des renseignements tels que la fréquence des allèles, la probabilité d'identité, l'hétérozygotie attendue et observée ainsi que les paires probables de parents-enfants. Ces logiciels permettent aussi de prendre en considération les génotypes incomplets à certains loci pour identifier les individus identiques.

Estimation de la taille de la population (N)

Une fois les génotypes identiques identifiés et les événements de recapture inclus dans l'historique de capture, il est possible d'estimer la taille de la population ainsi que sa densité. Il existe différentes approches pour estimer la taille et la densité de la population, selon que la population soit considérée comme ouverte ou fermée, ou que nous soyons intéressés par la densité seulement. Le programme CloseTest permet de vérifier si la population est fermée démographiquement en estimant le degré de violation de la fermeture démographique en plus de fournir des informations sur la nature de la violation et les périodes pour lesquelles cette condition est violée (Stanley et Richards, 2005). Ce programme permet de vérifier la fermeture géographique selon deux tests : 1) celui développé par Otis et coll. (1978) qui vérifie le degré de fermeture considérant qu'il n'y a

pas de variabilité dans la probabilité de capture entre les individus; 2) celui de Stanley et Burnham (1999) qui teste l'hypothèse nulle que la population soit fermée en considérant la variabilité dans la probabilité de capture entre les individus contre l'hypothèse alternative que la population soit ouverte selon le modèle de Jolly-Seber. Le test de Stanley et Burnham (1999) est le plus sensible puisqu'il considère la variation dans la probabilité de capture. Il permet aussi de déterminer la nature de la violation, à savoir si elle provient de l'ajout ou du retrait d'individus, et pendant quelle période la violation a été détectée (Stanley et Jon, 2005). Il existe plusieurs modèles permettant d'estimer la taille d'une population. La structure de base de tous ces modèles est la suivante :

$$\hat{N} = \frac{M}{\hat{p}^*} \quad (\text{Équation 3})$$

où $\hat{N}$ correspond au nombre d'individus, M au nombre total d'individus uniques marqués et $\hat{p}^*$ à la probabilité qu'un animal soit capturé au moins une fois. Cet estimateur sert de base aux modèles actuellement utilisés, que ce soit pour des populations fermées ou ouvertes. Les principaux modèles d'estimation de l'abondance pour des populations fermées sont présentés au tableau 2. Ces modèles permettent de considérer l'effet du temps, de la réponse comportementale des individus suite à une première capture (comportements « trappophile » ou « trappophobe ») et la variabilité dans la probabilité de capture entre les individus. Ces effets peuvent être considérés séparément ou en combinaison. La modélisation de ces effets limite les conséquences de ne pas respecter toutes les suppositions, telles que la probabilité de capture identique pour tous les individus et constante dans le temps, ainsi que la supposition que le marquage n'influence pas la probabilité de recapture.

Tableau 2. Description des modèles d'estimation de la taille d'une population fermée selon la prise en compte des effets du temps, de la réponse comportementale à une première capture et de la variation individuelle dans la probabilité de capture.

Modèle	Description
M_o	La probabilité de capture est constante dans le temps et indépendante des captures antérieures.
M_t	La probabilité de capture est variable dans le temps et indépendante des captures antérieures.
M_b	La probabilité de capture est influencée par la réponse comportementale à une première capture.
M_{tb}	La probabilité de capture est variable dans le temps et influencée par la réponse comportementale à une première capture.
M_h	La probabilité de capture est variable entre les individus et indépendante des captures antérieures.
M_{th}	La probabilité de capture est variable dans le temps, variable entre les individus et indépendante des captures antérieures.
M_{bh}	La probabilité de capture est influencée par la réponse comportementale à une première capture et variable entre les individus.
M_{tbh}	La probabilité de capture est variable dans le temps, variable entre les individus et influencée par la réponse comportementale à une première capture.

Plusieurs logiciels sont disponibles pour estimer la taille d'une population (voir le site Web <http://www.phidot.org/>). Les logiciels CAPTURE (Otis et coll., 1978; Rexstad et Burnham, 1992) et MARK (White et Burnham, 1999) sont sans aucun doute les plus utilisés. Le logiciel CAPTURE permet d'estimer la taille de la population avec la plupart des modèles, bien que certaines combinaisons présentées dans le tableau 2 ne soient pas disponibles. Le modèle le mieux supporté par les données empiriques est sélectionné à l'aide du critère d'information d'Akaike [AIC] (Akaike, 1974). Le programme MARK offre la possibilité de tester tous les modèles présentés au tableau 2 et d'estimer la population à l'aide d'une moyenne pondérée selon le poids de chacun des modèles testés. Toutefois, ces deux logiciels ne fournissent pas d'estimation pour les effets du temps, de la réponse comportementale et de la variation dans la probabilité de capture entre les individus. Par exemple, si le logiciel CAPTURE propose le modèle M_t comme meilleur modèle, il est seulement possible d'affirmer que le temps a un effet sur la probabilité de capture, mais il est impossible de déterminer l'importance de cet effet et sa nature (positif ou négatif).

Il arrive souvent que les suppositions de fermeture démographique et géographique soient difficiles, voire impossibles à rencontrer en milieu naturel. Les effets d'une violation de ces postulats seront discutés à la section intitulée « Sources de biais des méthodes de CMR par génotypage ». Dans le cas où la population ne peut être considérée comme fermée, il est possible d'avoir recours aux modèles de populations ouvertes de Jolly-Seber (Pollock et coll., 1990) ou de Comark-Jolly-Seber (Lebreton et coll., 1992). Les modèles de populations ouvertes requièrent les mêmes suppositions que ceux pour les populations fermées, à l'exception du critère de fermeture de la population. En plus du taux de survie (ϕ), de la probabilité de capture (p) et de la taille de la population (N) qui sont aussi estimés dans les modèles de populations fermées, on obtient une estimation du nombre de nouveaux individus entrant dans la population (B). Malgré l'abondante littérature sur les modèles de populations ouvertes (Pollock et coll., 1990; Lebreton et coll., 1992; Schwarz et Arnason, 1996), les modèles de populations fermées sont souvent préférés. Ces derniers sont considérés comme plus précis et comportent moins de biais que ceux pour les populations ouvertes, parce qu'ils tiennent compte de la variation dans la probabilité de capture dans le temps, des changements comportementaux et des différences individuelles

(voir Pledger, 2000, pour plus de détails à ce sujet). Ils donnent ainsi des estimations plus exactes et précises que les modèles de populations ouvertes, particulièrement pour les espèces pour lesquelles une grande variation comportementale est vraisemblable, comme chez les ursidés.

Estimation de la densité d'une population

Les données de CMR permettent aussi d'estimer la densité d'une population, à partir de l'estimation de l'abondance issue de l'historique de capture ou directement avec cet historique, sans passer par l'estimation de l'abondance. La première méthode est certainement la plus répandue et la plus intuitive. Pour obtenir une estimation de la densité D , on divise l'abondance N par la taille de l'aire d'étude A :

$$D = \frac{N}{A} \quad (\text{Équation 4})$$

Toutefois, pour estimer la densité, il faut connaître la superficie occupée par la population échantillonnée. La superficie de l'aire d'étude peut être estimée avec un polygone convexe entourant toutes les stations d'échantillonnage (A). Cependant, des individus dont seulement une partie du domaine vital chevauche l'aire d'étude pourraient être détectés aux stations d'échantillonnage situées en périphérie de l'aire d'étude, ce qui a pour effet de surestimer la densité puisque l'aire réellement couverte par les stations d'échantillonnage est sous-estimée (Otis et coll., 1978). Il a donc été proposé d'ajouter une bande de largeur W autour de chaque station pour déterminer une aire d'étude effective $\hat{A}$ (Dice, 1938). Plusieurs méthodes ont été développées pour déterminer la valeur de la bande entourant les stations (W). Il a été suggéré d'accorder à W la valeur moyenne de la distance maximale parcourue par chaque individu recapturé [MMDM] (Wilson et Anderson, 1985), ou la moitié de cette valeur [MMDM/2] (Otis et coll., 1978). Les mouvements mis en évidence par les événements de recapture sur une grille de stations de collecte de poils sont souvent sous-estimés puisqu'ils sont tronqués aux stations (Obbard et coll., 2010). Les valeurs de MMDM et MMDM/2 sont donc susceptibles de sous-estimer la superficie de l'aire d'étude. L'utilisation des distances de déplacement issues des données télémétriques, donc indépendantes de la position spatiale des stations, permet de corriger ce biais. Jett et

Nichols (1987) ont développé une méthode plus rigoureuse pour déterminer W , soit la moitié de la distance asymptotique entre les événements de capture et de recapture, ou $ARL/2$. Elle présume que plus le nombre de recaptures d'un individu est élevé, plus il y a de chances que la distance maximale entre deux sites de capture soit grande. Toutefois, cette distance maximale atteint un plateau correspondant à la capacité de l'individu à se déplacer (figure 3). La moitié de la distance maximale atteinte par le plateau correspond à $ARL/2$ (figure 3). Il est aussi possible, entre autres avec le logiciel DENSITY 4.4, d'ajuster manuellement la valeur de W . Par exemple, si nous avons des raisons biologiques de croire que W devrait être calculé différemment qu'en utilisant les méthodes précédentes, il est possible d'ajuster sa valeur. Bien que l'utilisation de l'aire d'étude effective $\hat{A}$ dans le calcul de la densité soit une méthode efficace pour contrer l'effet de bordure, ces méthodes sont encore imparfaites et elles ne sont pas nécessairement applicables à toutes les situations. En effet, l'obtention d'une valeur de MMDM ou MMDM/2 non biaisée nécessite l'utilisation de la télémétrie, ce qui augmente considérablement les coûts de l'étude. Quant à $ARL/2$, il faut obtenir suffisamment d'événements de recapture pour atteindre une asymptote, ce qui n'est généralement pas facile avec de grands mammifères élusifs comme les ursidés.

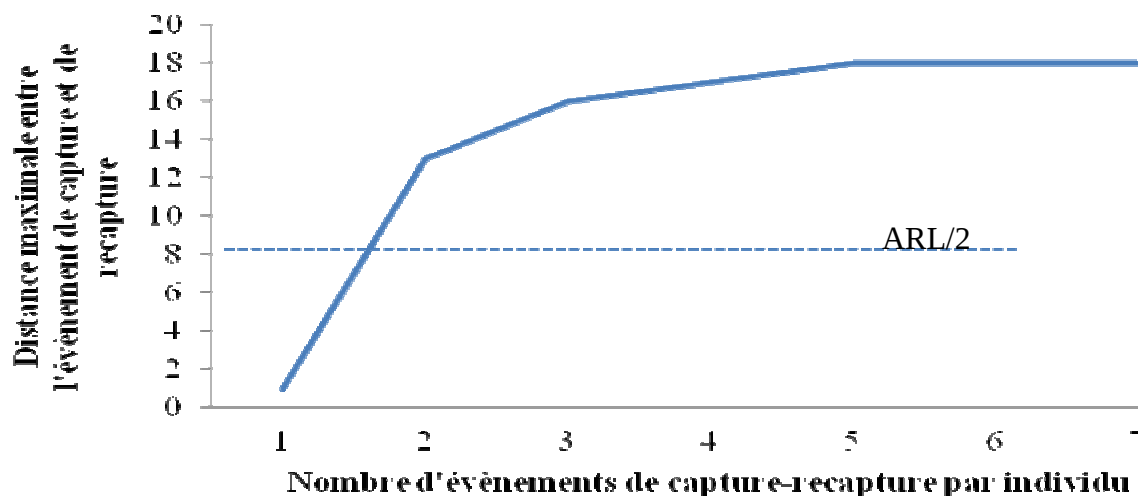


Figure 2. Représentation de la moitié de la distance asymptotique maximale entre un évènement de capture et un évènement de recapture.

D'autres méthodes ont récemment été mises au point pour estimer la densité sans recourir aux estimations de N et de $\hat{A}$, tout en incluant une dimension spatiale dans la modélisation de la probabilité de capture. Parmi ces méthodes, on retrouve la prédiction inverse pour des

captures-recaptures spatialement explicites [IP SECR] (Efford et coll., 2004) et la maximisation de la vraisemblance des captures-recaptures spatialement explicites [ML SECR] (Efford, 2004). Les deux techniques se basent sur l'estimation de trois paramètres : la densité D , la fonction de probabilité de détection en fonction de la distance à une station d'échantillonnage $g(d)$ et l'étendue spatiale sur laquelle la probabilité de détection diminue $s(d)$. Alors que l'IP SECR se base sur la simulation et la prédiction inverse des paramètres, la ML SECR se base sur la maximisation de la vraisemblance (Efford et coll., 2009). La méthode de la ML SECR a remplacé celle de la prédiction inverse puisqu'elle permet une plus grande flexibilité dans la sélection des modèles. Les deux techniques sont disponibles dans le logiciel DENSITY 4.4 (Efford, 2009) et la ML SECR est aussi disponible dans l'environnement de travail R (R Development Core Team, 2011).

L'approche ML SECR permet de modéliser les effets du temps, de la réponse comportementale à une première capture et de la variation individuelle dans la probabilité de capture (tableau 3). Le modèle le plus simple de la ML SECR utilise deux sous-modèles : un sous-modèle représentant la distribution spatiale des individus, au moyen d'une grille de points (x_i, y_i) représentant le centre hypothétique des domaines vitaux des individus, et un autre sous-modèle représentant une fonction de probabilité de capture à une station $g(d)$ en fonction de la distance entre la station et le centre hypothétique des domaines vitaux (x_i, y_i) . Les coordonnées du centre des vrais domaines vitaux (x_i, y_i) sont inconnues, mais on suppose qu'elles suivent une distribution de Poisson sur la longitude et la latitude. Ces équations sont résolues en estimant simultanément D , $g(d)$ et $s(d)$ par la maximisation de la vraisemblance. On doit préciser la distribution que suit la fonction de détection $g(d)$, souvent présumée comme demi-normale, mais qui peut aussi être de forme exponentielle négative ou aléatoire. La maximisation de la vraisemblance peut se réaliser en utilisant la vraisemblance complète ou conditionnelle. Contrairement aux logiciels CAPTURE et MARK, il n'existe pas d'analogue au modèle M_t (effet du temps, plus spécifiquement l'effet de l'occasion de capture sur la probabilité de capture) dans le logiciel DENSITY 4.4. Une fois les différentes combinaisons des effets du temps, de la réponse comportementale et de la variation individuelle prises en compte dans la

probabilité de capture, il est possible de comparer les modèles par leur AIC ou AICc (AIC corrigé pour des échantillons de faible taille et qui pénalise davantage pour des paramètres additionnels). Il est aussi possible de calculer une moyenne pondérée à partir des poids d'AIC.

La ML SECR offre plusieurs avantages par rapport à la méthode de dérivation de la densité avec des estimations de la taille de population et de la superficie effective de l'aire d'étude. D'abord, il n'est pas nécessaire d'estimer N et $\hat{A}$, paramètres pour lesquels une définition biologiquement significative est difficile. La valeur de N obtenue par estimation représente en fait une sous-population « capturable » par les stations d'échantillonnage provenant d'une plus grande population de taille inconnue qui s'étend au-delà de la zone de capture. La relation entre cette estimation de N et la taille réelle de la population reste mal définie (Efford, 2009). Il est donc préférable d'aborder l'estimation de la densité avec une méthode ne requérant pas l'estimation de N et de $\hat{A}$. De plus, la ML SECR permet d'évaluer la part de la variation dans la probabilité de capture entre les individus qui est causée par la répartition spatiale des stations d'échantillonnage, ce que ne permettent pas les méthodes d'estimation de la taille des populations. La modélisation s'applique donc à n'importe quel type de disposition spatiale des stations. Finalement, la méthode est jugée robuste au choix de la fonction de détection $g(d)$ (Efford, 2009). Il est à noter que la précision des estimations issues de la maximisation de la vraisemblance augmente avec le nombre de recaptures (Efford et coll., 2004).

Le logiciel DENSITY 4.4 fournit donc à la fois une plateforme pour réaliser la maximisation de la vraisemblance pour des captures-recaptures spatialement explicites et l'estimation de la taille et de la densité de population à partir de modèles de populations fermées et d'aires effectives. L'estimation de l'abondance de la population se fait à l'aide du programme CAPTURE qui est intégré à une fenêtre du programme DENSITY 4.4. L'environnement de travail R fournit aussi une plateforme intéressante pour réaliser des ML SECR. Le codage est simple et expliqué dans la documentation de l'extension *secur* (Efford, 2011).

Tableau 3. Modèles d'estimation de la densité pour les populations fermées selon la maximisation de la vraisemblance à partir d'un historique de capture-recapture spatialement explicite. Les modèles incluent un ou plusieurs effets suivants : le temps (t), la réponse comportementale à une première capture (b) et la variation individuelle dans la probabilité de capture (h).

Modèle	Description
$g(.)s(.)$	La fonction de probabilité de capture et l'étendue spatiale sur laquelle la probabilité de capture diminue sont constantes.
$g(.)s(h)$	La fonction de probabilité de capture est constante, mais l'étendue spatiale sur laquelle la probabilité de capture diminue est variable entre les individus.
$g(t)s(.)$	La fonction de probabilité de capture est influencée par le temps, mais l'étendue spatiale sur laquelle la probabilité de capture diminue est constante.
$g(t)s(h)$	La fonction de probabilité de capture est influencée par le temps et l'étendue spatiale sur laquelle la probabilité de capture diminue est variable entre les individus.
$g(b)s(.)$	La fonction de probabilité de capture est influencée par la réponse comportementale après une première capture, mais l'étendue spatiale sur laquelle la probabilité de capture diminue est constante.
$g(b)s(h)$	La fonction de probabilité de capture est influencée par la réponse comportementale après une première capture et l'étendue spatiale sur laquelle la probabilité de capture diminue est variable entre les individus.
$g(h)s(.)$	La fonction de probabilité de capture est variable entre les individus, mais l'étendue spatiale sur laquelle la probabilité de capture diminue est constante.
$g(h)s(h)$	La fonction de probabilité de capture et l'étendue spatiale sur laquelle la probabilité de capture diminue sont variables entre les individus.
$g(tb)s(.)$	La fonction de probabilité de capture est influencée par le temps, mais la réponse comportementale après une première capture et l'étendue spatiale sur laquelle la probabilité de capture diminue est constante.
$g(tb)s(h)$	La fonction de probabilité de capture est influencée par le temps et la réponse comportementale après une première capture et l'étendue spatiale sur laquelle la probabilité de capture diminue est variable entre les individus.

Tableau 3 (suite)

Modèle	Description
$g(th)s(.)$	La fonction de probabilité de capture est influencée par le temps et variable entre les individus, mais l'étendue spatiale sur laquelle la probabilité de capture diminue est constante.
$g(th)s(h)$	La fonction de probabilité de capture est influencée par le temps et variable entre les individus, mais l'étendue spatiale sur laquelle la probabilité de capture diminue est variable entre les individus.
$g(bh)s(.)$	La fonction de probabilité de capture est influencée par la réponse comportementale après une première capture et variable entre les individus, mais l'étendue spatiale sur laquelle la probabilité de capture diminue est constante.
$g(bh)s(h)$	La fonction de probabilité de capture est influencée par la réponse comportementale après une première capture et variable entre les individus et l'étendue spatiale sur laquelle la probabilité de capture diminue est variable entre les individus.
$g(tbh)s(.)$	La fonction de probabilité de capture est influencée par le temps, la réponse comportementale après une première capture et variable entre les individus et l'étendue spatiale sur laquelle la probabilité de capture diminue est constante.
$g(tbh)s(h)$	La fonction de probabilité de capture est influencée par le temps, la réponse comportementale après une première capture et variable entre les individus et l'étendue spatiale sur laquelle la probabilité de capture diminue est variable entre les individus.

Sources de biais des méthodes de CMR par génotypage

Les modalités de gestion de la faune exploitée par la chasse et le piégeage dépendent des estimations de la taille et de la densité des populations. Ces estimations doivent être assez précises et non biaisées. La méthode de CMR par génotypage offre plusieurs avantages, bien qu'elle puisse être sujette à des biais importants causés par le design expérimental et les choix statistiques. Selon Proctor et coll. (2010), les estimations de population peuvent être biaisées lorsqu'il y a 1) violation des suppositions de fermeture géographique et démographique, 2) déséquilibre entre la superficie de l'aire d'étude, l'intensité de l'échantillonnage, le nombre de captures, la probabilité de capture et la puissance statistique associée et 3) variation interindividuelle dans la probabilité de capture. Les autres sources de biais spécifiques aux études de CMR par génotypage sont 4) la façon de considérer la capture d'un même individu à différents endroits lors d'une même session d'échantillonnage (Lukacs et Burnham, 2005) et les visites de plusieurs individus à une station lors d'une même session d'échantillonnage (Gagneux et coll., 1997).

1) Violation des suppositions de fermeture géographique et démographique

Comme mentionné précédemment, les modèles de CMR pour des populations dites fermées sont reconnus comme étant plus précis et moins biaisés. Cependant, les suppositions de fermeture géographique et démographique des modèles de CMR peuvent être violées dans des populations naturelles. Or, une telle violation peut grandement biaiser les estimations de population si elle n'est pas considérée dans le modèle de CMR (Boulanger et coll., 2002). Il est cependant possible de limiter la probabilité de violer ces suppositions avec un design expérimental et une approche de modélisation adéquats.

La fermeture démographique est difficile, voire impossible à contrôler complètement. La violation de cette supposition peut induire des biais importants dans les estimations de la taille et de la densité d'une population. En effet, la mortalité d'individus suite à leur capture surestime la taille réelle de la population, alors que la naissance d'individus pendant la période d'échantillonnage sous-estime la taille réelle de la population. Il est possible de

limiter la violation de la fermeture démographique à l'aide du design expérimental. Par exemple, pour les ursidés, il est préférable de réaliser l'échantillonnage pendant l'été pour éviter qu'il y ait des naissances et minimiser les risques de mortalité (p. ex., par la chasse). Le logiciel CloseTest permet de vérifier le respect de cette condition.

En ce qui concerne la supposition de fermeture géographique, le fait d'échantillonner des individus qui n'appartiennent pas à l'aire d'étude a pour effet de surestimer la taille de la population. À l'inverse, ne pas capturer des individus qui appartiennent à l'aire d'étude, mais qui ne l'ont pas fréquentée pendant la période d'échantillonnage sous-estime la taille de la population. On peut tenter de limiter les échanges en délimitant l'aire d'étude en fonction de barrières naturelles et anthropiques. Proctor et coll. (2010) suggèrent que l'utilisation de barrières anthropiques importantes, comme les autoroutes et les villes, serait très efficace pour fermer géographiquement une population dans le cadre d'une étude de CMR (mais voir Poole et coll., 2001).

La superficie de la grille d'échantillonnage par rapport à la capacité de déplacement de l'espèce est un autre critère important. En effet, une grille relativement petite par rapport à la taille moyenne des domaines vitaux peut induire un effet de bordure important (Boulanger et coll., 2002). En augmentant la superficie de l'aire d'étude par rapport à la superficie des domaines vitaux, la proportion d'individus susceptibles de quitter l'aire d'étude à un moment ou un autre de la période d'échantillonnage diminue, ce qui permet de réduire les biais liés à l'effet de bordure (Proctor et coll., 2010). De plus, il est possible de réaliser l'inventaire pendant la période où les animaux sont le moins mobiles afin de limiter la violation de la supposition de fermeture géographique. Si des individus ont été munis d'un collier émetteur dans l'aire d'étude, il est possible de vérifier la fidélité à l'aire d'étude et de corriger les estimations de population en utilisant la proportion de temps que les individus passent à l'intérieur de celle-ci (Boulanger et coll., 2004).

2) Équilibre entre les paramètres de capture et la puissance statistique

Comme les ursidés en général, l'ours noir est une espèce évasive pour laquelle les taux de capture et de recapture sont généralement faibles, ce qui rend les estimations de population imprécises. Selon White et coll. (1982), la probabilité de capture devrait être supérieure à 0,20 pour des tailles de population estimées entre 100 et 200 individus, alors qu'elle devrait être supérieure à 0,30 pour des tailles de population inférieures à 100 individus, afin que la procédure de sélection de modèles soit fiable. Cependant, pour respecter les contraintes en ressources humaines et financières, un compromis entre la superficie de l'aire d'étude et l'intensité d'échantillonnage est nécessaire. Le design expérimental peut toutefois être adapté de façon à maximiser les probabilités de capture et de recapture. Par exemple, le positionnement des parcelles d'échantillonnage dans les meilleurs habitats disponibles pour l'ours permet d'obtenir des taux de capture et de recapture plus élevés (p. ex., Woods et coll., 1999; Boulanger et coll., 2004; Dreher et coll., 2007). À cet effet, il faut savoir que les forêts riches en feuillus qui produisent des fruits secs indéhiscents (frêne et chêne) sont considérées comme des habitats de choix pour l'ours noir dû à la qualité nutritionnelle de ces ressources (Samson et Huot, 1994; Samson, 1996). Les forêts perturbées de façon naturelle ou par des activités anthropiques, comme les coupes forestières et les brûlis en régénération, constituent aussi des habitats favorables et fortement utilisés par l'ours (Boileau et coll., 1994; Samson et Huot, 1994; Jolicoeur, 2004; Leblanc et Samson, 2006). Les milieux non forestiers sont généralement évités, bien que les champs agricoles soient parfois visités pour l'alimentation (Leblanc et Samson, 2006).

Enfin, la superficie des parcelles doit être inférieure à celle des domaines vitaux du sexe le moins mobile pour assurer l'accessibilité aux stations d'échantillonnage à tous les individus présents dans l'aire d'étude (Proctor et coll., 2010). Le déplacement des stations d'échantillonnage entre les sessions de capture permet de capturer plus de femelles et d'éliminer la réponse comportementale que peuvent développer certains individus à des stations fixes (Proctor et coll., 2010).

3) Variation interindividuelle dans la probabilité de capture

Plusieurs études ont mis en évidence la présence de variation dans la probabilité de capture entre les individus chez les ursidés (p. ex., Boulanger et coll., 2002; Gardner et coll., 2010; Proctor et coll., 2010). Cette variation peut être prise en compte dans les logiciels comme MARK, CAPTURE ou R. Cependant, lorsque le taux de capture est relativement faible, la variation peut paraître inexistante, même si elle est présente en réalité. Par exemple, un taux de capture inférieur à 0,20 chez les ours signifie que la majorité d'entre eux ne sont capturés qu'une seule fois et il est impossible de quantifier la variation interindividuelle avec un seul événement de capture. Les estimations de population peuvent donc être fortement biaisées, particulièrement pour de petites populations (Proctor et coll., 2010).

4) Captures multiples d'un individu lors d'une session d'échantillonnage et capture de plusieurs individus à une même station

Les techniques de CMR par génotypage permettent aux individus de poursuivre leurs déplacements après l'évènement de capture. De plus, la station d'échantillonnage demeure souvent disponible pour d'autres événements de capture suite à une première capture. Ainsi, il est possible de capturer un même individu à plusieurs stations lors d'une même session d'échantillonnage. Ces événements de capture multiples ne sont pas en soit une source de biais, mais la façon dont ils sont traités statistiquement peut l'être. La plupart des études regroupent ces événements de capture « simultanés » dans un seul événement de capture pour la session d'échantillonnage, perdant ainsi de l'information précieuse. Évidemment, augmenter la fréquence des visites aux stations d'échantillonnage peut réduire la probabilité d'observer cette situation. Cependant, ceci requiert une charge de travail supplémentaire et augmente les coûts associés au travail de terrain. Les développements récents dans la modélisation des historiques de CMR permettent maintenant de prendre en considération ces événements de capture multiples en basant l'estimation de population sur un échantillonnage spatialement explicite (Lukacs et Burnham, 2005). Par exemple, le logiciel DENSITY 4.4 permet de spécifier la possibilité que les stations d'échantillonnage soient visitées par plus d'un individu par occasion de capture et le format de l'historique de capture permet de conserver les événements de

capture simultanés pour un même individu (option du logiciel : type de trappe « proximity »).

Sources d'erreurs des analyses génétiques et solutions

Bien que la méthode de CMR par génotypage des poils permette de limiter les biais liés, par exemple, à la fermeture géographique ou démographique, elle introduit toutefois de nouvelles sources de biais et d'erreurs liées aux analyses génétiques. Bellemain et coll. (2005) ont testé quatre différentes méthodes d'estimation de la taille d'une population d'ours bruns à l'aide de données génétiques provenant de la récolte de fèces, soit l'analyse de raréfaction, un estimateur de CMR basé sur les données génétiques, un estimateur de CMR basé sur des données génétiques et d'observation, et une analyse de la survie des animaux portant des colliers émetteurs. Ils ont conclu que la méthode d'estimation se basant exclusivement sur les données génétiques était la plus fiable, tout en étant la plus sensible aux sources d'erreur. Il existe plusieurs types d'erreurs reliées au génotypage : l'échec de l'amplification de certains allèles (*allelic dropout*), les faux allèles (*false allele*), l'effet d'ombre (*shadow effect*), les erreurs de lecture des bandes et les erreurs de transcription (Taberlet et coll., 1996; Gagneux et coll., 1997; Goossens et coll., 1998; Taberlet et coll., 1999; Mills et coll., 2000; Waits et Leberg, 2000; Wright et coll., 2009).

Les erreurs les plus fréquentes sont celles causées par l'échec de l'amplification d'allèles. Celui-ci survient quand les allèles d'hétérozygotes ne sont pas amplifiés lors de la PCR, menant à l'identification erronée d'individus homozygotes (Gagneux et coll., 1997). Ce problème peut être causé par une quantité trop faible d'ADN utilisée pour l'analyse ou une mutation au site de l'allèle qui empêche une amplification juste. Les estimations de population peuvent donc être surestimées en ajoutant des individus inexistant dans la population réelle (Wright et coll., 2009). Les faux allèles surviennent lorsqu'une erreur s'introduit dans le génotype au moment de l'amplification, créant ainsi un « faux » génotype et introduisant un « faux » individu dans la population, lequel ne sera jamais recapturé, à moins que l'erreur de génotypage se répète. Cela a tendance à surestimer la population réelle (Taberlet et Luikart 1999; Waits et Leberg, 2000). L'effet d'ombre est une

erreur d'identification des individus due à un nombre insuffisant de marqueurs microsatellites utilisés, une trop faible variabilité des microsatellites entre les individus de la population ou une mutation lors de l'amplification d'un allèle. Il peut en résulter une sous-estimation de la population puisque des individus différents se voient attribuer des génotypes identiques (Mills et coll., 2000; Waits et Leberg, 2000).

Les erreurs de génotypage peuvent être corrigées, ou du moins réduites, par des méthodes rigoureuses en laboratoire (Paetkau, 2003). Entre autres, la réplication des analyses génétiques pour chaque échantillon s'avère une méthode efficace pour détecter les erreurs de génotypage. Il existe trois approches pour la réplication des analyses génétiques. La première, nommée approche multitubes, consiste à analyser tous les échantillons plusieurs fois et à déterminer les génotypes en se basant sur les patrons les plus fréquents. Navidi et coll. (1992) ainsi que Taberlet et coll. (1996) ont proposé l'adoption de l'approche multitubes lorsque l'ADN est disponible en quantité suffisante. Cependant, la réanalyse de certains échantillons demande une dilution plus importante de l'ADN total disponible, diminuant par conséquent la quantité d'ADN disponible pour chaque analyse. De plus, les coûts associés aux analyses génétiques montent en flèche avec l'ajout de plusieurs tubes pour chaque échantillon (Waits et Paetkau, 2005). Enfin, bien qu'une erreur de génotypage survienne dans moins de 1 % des échantillons avec l'approche multitubes, Waits et Leberg (2000) ont tout de même observé une surestimation des populations de l'ordre de 10 à 30 %. Il est possible de déterminer, lors d'une étude préliminaire, le nombre de tubes requis pour atteindre un niveau d'erreur acceptable tout en limitant les coûts (Goossens et coll., 1998).

La deuxième approche de réplication consiste à réanalyser seulement les échantillons qui ont des génotypes identiques. Puisque les erreurs de génotypage ont très peu de chances de se reproduire exactement aux mêmes loci pour les mêmes échantillons, cette analyse permet de différencier les échantillons contenant réellement des génotypes identiques des paires issues d'une erreur de génotypage (Woods et coll., 1999; Waits et Paetkau, 2005). Avec cette stratégie, il est possible d'utiliser près de 30 % de la quantité totale d'ADN disponible pour la première analyse, limitant ainsi les erreurs causées par de trop petites

quantités d'ADN tout en limitant les coûts (Waits et Paetkau, 2005). Finalement, la troisième approche de réplication repose sur l'identification des loci hétérozygotes pour des génotypes identifiés comme semblables parmi la population échantillonnées et la réplication de ces loci seulement. Toutefois, il semble que cette méthode ne soit pas suffisamment efficace pour abaisser le taux d'erreurs de génotypage à moins de 5 % (Roon et coll., 2005). L'approche de réplication ciblée pour les génotypes identiques semble donc présenter le meilleur rapport coûts-bénéfices (Woods et coll., 1999; Waits et Paetkau, 2005).

Les faux allèles surviennent lorsque des bandes d'artéfacts de même intensité que celles des vrais allèles apparaissent sur les gels en raison d'un glissement lors de la séquence d'amplification (Taberlet et coll., 1996). L'utilisation de plus grandes quantités d'ADN pour chaque échantillon semble diminuer l'intensité des bandes d'artéfacts par rapport aux bandes réelles et ainsi diminuer les risques d'erreur lors de l'attribution des génotypes. De plus, les microsatellites tri- ou tétranucléotides sont moins susceptibles aux glissements de bandes d'artéfacts (Goossens et coll., 1998). L'effet d'ombre se produit lorsqu'en augmentant le nombre de marqueurs microsatellites, on augmente la probabilité d'identifier des génotypes différents à partir d'échantillons d'ADN provenant en réalité du même individu. Logiquement, on pourrait croire qu'augmenter le nombre de marqueurs microsatellites augmente notre capacité à discriminer deux individus différents (McKelvey et Schwartz, 2004), mais Waits et Leberg (2000) ont montré que plus il y a de marqueurs utilisés, plus la probabilité qu'il y ait une erreur de génotypage sur au moins un de ces marqueurs augmente. Il est préférable de choisir un nombre réduit de microsatellites hautement polymorphiques que de choisir un grand nombre de microsatellites peu ou moyennement polymorphiques. On réduit alors les erreurs de génotypage tout en maximisant l'unicité des combinaisons alléliques (Waits et Leberg, 2000). Il est cependant difficile de prévoir comment les microsatellites performeront avec les individus capturés lors de l'étude. Il est possible de déterminer, à l'aide d'un sous-échantillon de poils provenant d'individus fréquentant le site d'étude, le nombre optimal de marqueurs microsatellites à utiliser pour minimiser les erreurs de génotypage et maximiser la différenciation individuelle (Mills et coll., 2000; Paetkau, 2003). Waits et Paetkau (2005)

croient qu'il est prudent d'utiliser un marqueur de plus que recommandé jusqu'à ce que le nombre d'échantillons recueillis soit suffisant pour démontrer que l'abandon d'un microsatellite ne fait pas augmenter la proportion de génotypes semblables à un locus.

La quantité d'ADN disponible dans les échantillons est sans aucun doute un facteur limitant important pour les analyses génétiques à l'aide de poils (Taberlet et coll., 1996; Gagneux et coll., 1997; Taberlet et coll., 1997). Il est possible que la quantité d'ADN soit suffisante pour que la séquence d'amplifications fonctionne, mais qu'elle soit insuffisante pour différencier les loci homozygotes des loci hétérozygotes. En effet, Taberlet et coll. (1996) ont démontré que les échecs d'amplification d'allèles sont plus fréquents pour des quantités d'ADN inférieures à 56 pg (résultat d'amplification positif, mais identification incorrecte). Cette conclusion est aussi supportée par les résultats de Gagneux et coll. (1997) [seuil à 0,5 ng d'ADN / 10 µl de solution]. Il est donc possible de diminuer les erreurs liées à la délétion d'allèles en augmentant la quantité d'ADN traitée à chaque analyse. Il restera impossible d'éliminer complètement cette source d'erreur à cause de la qualité et de la faible quantité de matériel généralement recueilli sur le terrain (Goossens et coll., 1998). Toutefois, l'utilisation de poils ayant obligatoirement leurs follicules intacts, donc fournissant plus d'ADN, favorise la diminution des erreurs de génotypage. En utilisant des poils de marmottes (*Marmota marmota*) avec follicules, Goossens et coll. (1998) ont démontré que le taux d'erreur de génotypage pouvait passer de 14,0 % pour des échantillons avec un seul poil à 0,29 % pour des échantillons avec dix poils. La plupart des études sur les ursidés limitent d'ailleurs les analyses génétiques aux échantillons contenant cinq poils avec follicules ou plus (Settlage et coll., 2008). Toutefois, abandonner les échantillons ne contenant pas suffisamment de poils pourrait, selon Lukacs et Burnham (2005), introduire une source de variation supplémentaire dans la probabilité de capture. En effet, il pourrait exister une différence individuelle dans le taux de perte de poils, certains individus ayant tendance à perdre moins de poils. Ainsi, éliminer les échantillons contenant peu de poils pourrait réduire grandement les chances que ces individus soient identifiés. Lukacs et Burnham (2005) soutiennent qu'il est préférable de conserver ces petits échantillons malgré le risque d'erreur de génotypage associé afin d'échantillonner de façon uniforme les individus d'une population. Enfin, puisque les échantillons contenant peu

d'ADN sont plus sujets à la contamination, le laboratoire responsable des analyses devrait faire attention pour éviter toute contamination croisée lors de l'analyse (Waits et Paetkau, 2005).

Comme le soulignent plusieurs études, les biais en apparence mineurs peuvent mener à des erreurs importantes d'estimation de la taille des populations (Taberlet et coll., 1996; Mills et coll., 2000; Waits et coll., 2001; Boulanger et coll., 2002; Creel et coll., 2003; Paetkau, 2003; McKelvey et Schwartz, 2004). Par exemple, Creel et coll. (2003) ont obtenu une estimation de population 5,5 fois supérieure à la taille réelle d'une population de loup gris à cause d'erreurs de génotypage. La plupart des sources d'erreur de génotypage augmentent le nombre de génotypes uniques identifiés et la probabilité qu'un individu soit capturé au moins une fois est sous-estimée. Il est donc d'une grande importance d'effectuer des tests préliminaires pour détecter et quantifier chacune des sources d'erreur, et ce, afin d'élaborer un design expérimental qui réduit au minimum les sources d'erreur (Goossens et coll., 1998).

Bien qu'il soit possible de déterminer et de quantifier les sources d'erreur de génotypage, il demeure toutefois difficile d'établir un protocole généralement applicable qui permettrait de minimiser celles-ci, puisque la quantité et la qualité d'ADN varient grandement d'une étude à l'autre. Aussi, les erreurs potentielles discutées précédemment sont difficilement décelables sans avoir effectué un test spécifique. McKelvey et Schwartz (2004) ont proposé deux tests relativement simples pour quantifier les erreurs de génotypage. Le premier s'attarde à la surabondance de génotypes dits uniques, puisque les erreurs de génotypage provoquent généralement l'addition de génotypes fictifs dans la population. Le second test évalue le taux auquel de nouveaux individus sont découverts à mesure que l'on ajoute des loci dans la reconnaissance génotypique. La combinaison de ces deux tests donne une idée générale de la qualité des échantillons pour la génétique (ces tests sont réalisables à l'aide du logiciel Dropout, disponible gratuitement en ligne [McKelvey et Schwartz, 2005]). Il est aussi possible de tester des échantillons dont la provenance est connue et qui passent par les mêmes étapes que les autres échantillons pour estimer le taux total d'erreurs accumulées lors des procédures (Bonin et coll., 2004). Taberlet et Luikart (1999) ont suggéré d'évaluer,

préalablement à une étude, la probabilité d'identité avec des données génétiques issues d'un pré-échantillonnage, de déterminer un taux d'erreur satisfaisant pour la question biologique posée, et, finalement, de vérifier s'il est possible d'atteindre ce taux d'erreur avec des manipulations en laboratoire. Enfin, le niveau de confiance désiré devrait être déterminé en fonction de la conséquence des biais et erreurs dans l'estimation des tailles de population et de la probabilité de capturer des animaux apparentés (Waits et coll., 2001).

Des modèles permettant d'intégrer les erreurs de génotypage dans l'estimation des tailles de population ont récemment été développés. Lukacs et Burnham (2005) ont proposé un modèle d'estimation de la taille d'une population fermée pour laquelle les données de CMR sont issues d'au moins trois sessions d'échantillonnage. Toutefois, cette méthode suppose que les erreurs de génotypage n'entraînent jamais l'obtention d'un génotype identique à celui d'un individu présent au sein de la population et que les erreurs de génotypage ne mènent jamais à l'obtention d'un génotype identique. Selon eux, les génotypes identiques proviennent donc d'un même individu et ne peuvent être issus d'une erreur de génotypage. Cette supposition peut toutefois être violée sachant que l'effet d'ombre est une source d'erreur possible. De la même façon, la méthode développée par Wright et coll. (2009) pour intégrer les erreurs d'identification des génotypes aux estimations d'abondance présume que la délétion d'allèle est la seule cause possible d'erreur de génotypage. Knapp et coll. (2009) ont développé un autre type de modélisation permettant d'assouplir la supposition selon laquelle les erreurs de génotypage ne font qu'ajouter des individus différents à la population, en plus de permettre la modélisation de données de CMR issues d'une ou de deux sessions d'échantillonnage. Cette méthode utilise les fréquences alléliques et les taux estimés d'erreurs de génotypage pour quantifier la variabilité dans les estimations de taille de population causée par l'incertitude des appariements de génotype. Plus particulièrement, pour chaque paire de génotypes obtenus (identiques ou non), la probabilité que cette paire de génotypes provienne d'un même individu est calculée. Puis, une « pseudo-matrice » d'historiques de capture qui peut être gérée par les logiciels d'estimation de paramètres de populations (p. ex., avec le logiciel MARK) est construite.

La prise en compte des erreurs d'identification des génotypes permet de réduire les efforts en laboratoire visant à faire disparaître complètement ces erreurs, comme l'approche multitubes. Au lieu de faire plusieurs réplicats avec peu d'ADN pour chaque échantillon, il est possible de se limiter à la première amplification réussie. Ainsi, on limite les coûts associés aux analyses génétiques. La modélisation de Knapp et coll. (2009) permet d'utiliser moins de marqueurs microsatellites puisqu'elle considère l'existence de l'effet d'ombre. Finalement, les erreurs causées par la réduction du nombre d'amplifications pour chaque échantillon et par un nombre réduit de microsatellites sur les estimations de la taille de population sont corrigées. Le logiciel MARK propose une série de modèles de populations fermées permettant d'intégrer à la fois la variation individuelle dans la probabilité de capture et le risque de mauvaise identification des génotypes. L'incorporation de ces deux types de variable dans un même modèle mène souvent à des résultats aberrants. Des erreurs d'identification de génotypes devraient, en général, surestimer le nombre d'individus dans la population, ce qui se reflète par un historique de capture où plusieurs individus ne sont capturés qu'une seule fois. Ceci a pour effet de faire augmenter la variation individuelle dans la probabilité de capture. Conséquemment, la modélisation est incapable de distinguer correctement les deux sources de variation, ce qui mène à des résultats biaisés (voir le fichier d'aide sur les modèles de populations fermées avec variation interindividuelle et mauvaise identification dans le logiciel MARK). Puisqu'il semble difficile d'isoler les sources d'erreurs génotypiques et qu'un protocole de laboratoire rigoureux permet de réduire considérablement les erreurs d'identification des génotypes, la limitation de ces erreurs avant les analyses démographiques semble un choix judicieux par rapport à l'intégration de leur probabilité dans les analyses. Bien que les méthodes d'analyse génétique engendrent de nouvelles sources d'erreurs dans les études de CMR non invasives, il demeure que cette technique est prometteuse dans l'étude des populations d'animaux rares, élusifs ou difficilement « capturables » ou manipulables, comme les ursidés (Waits et Paetkau, 2005).

4. INVENTAIRES RÉALISÉS EN ABITIBI-TÉMISCAMINGUE (2001-2003)

4.1. Description des études

Mise en contexte

Face à l'augmentation du nombre d'ours noirs récoltés par la chasse et le piégeage, de même qu'à une hausse probable du nombre de signalements d'ours importuns dans certaines régions du Québec, le Ministère a entrepris des études visant à élaborer une méthode d'inventaire abordable, applicable sur de grands territoires (zones de chasse) et permettant un suivi à long terme des populations. La méthode de CMR par génotypage des poils semblait une approche prometteuse. Les premiers travaux ont été réalisés lors de l'été 2000 dans la réserve des Laurentides. Ceux-ci visaient principalement à tester l'efficacité d'une technique de récolte de poils. Des analyses génétiques ont été réalisées sur les échantillons récoltés, mais aucune analyse démographique n'a été effectuée puisque ce n'était pas l'objectif (Lemieux et Desrosiers, 2002). La méthode a par la suite continué son processus d'élaboration lors des étés 2001 à 2003 dans la région de l'Abitibi-Témiscamingue (Courtois et coll., 2004). Les objectifs spécifiques étaient :

1. Récolter des échantillons de poils suffisamment nombreux et répartis adéquatement dans les territoires à l'étude;
2. Réaliser les analyses génétiques permettant la reconnaissance individuelle des individus capturés;
3. Estimer la taille des populations dans les aires d'étude;
4. Évaluer l'applicabilité de la technique pour estimer les populations d'ours noirs du Québec à l'échelle des zones de chasse;
5. Suggérer une méthode d'inventaire de l'ours noir applicable aux zones de chasse ou déterminer les connaissances manquantes pour en développer une.

Dans le cas de l'étude de 2001 en Abitibi-Témiscamingue, des analyses génétiques et démographiques ont été réalisées et les résultats ont été présentés dans le rapport d'étape de Courtois et coll. (2004). Nous avons refait les analyses des données issues de l'inventaire

de 2001 et avons comparé nos résultats à ceux de Courtois et coll. (2004). Nous avons également réalisé, lorsque possible, les analyses des données issues des inventaires de 2002 et 2003 pour lesquelles seulement les analyses génétiques avaient été effectuées.

Sites d'étude

Les sites d'étude des étés 2001 à 2003 étaient localisés dans la région de l'Abitibi-Témiscamingue, soit dans la zone de chasse 13. Les secteurs d'inventaire chevauchaient trois domaines bioclimatiques, soit le domaine de la sapinière à bouleau jaune, le domaine de la sapinière à bouleau blanc et celui de l'érablière à bouleau jaune. La faible densité des stations en 2003 s'explique par le recours à une disposition en grappe, avec l'objectif de couvrir un plus grand territoire, tout en maximisant l'accessibilité des stations pour tous les individus présents dans les grappes. Plus spécifiquement, l'inventaire a été réalisé autour de Rouyn-Noranda à l'été 2001 (R2001). À l'été 2002, il a été réalisé dans trois sites d'étude séparés, soit Roger-Rapide 7 (R2002), Senneterre (S2002) et Témiscamingue (T2002). À l'été 2003, les sites de Rouyn-Noranda (R2003-E et R2003-O), Témiscamingue (T2003), Senneterre (S2003) et Ville-Marie (V2003) ont été inventoriés. Le tableau 4 résume les informations relatives au nombre et à la densité de stations utilisées dans les différents projets. Les cartes montrant la disposition des stations sont présentées à l'annexe 2.

Dans le cas de R2001, deux stations d'échantillonnage ont été positionnées par parcelle de 40 km² puisque la superficie moyenne des domaines vitaux des femelles dans cette région était estimée à 20 km². Ceci devait permettre à tous les individus présents dans les grappes d'avoir accès à au moins une station. Dans tous les cas, la localisation exacte des stations a été déterminée en fonction de l'accessibilité par une route, un sentier ou une voie navigable.

Tableau 4. Nombre et densité (/20 km²) de stations d'échantillonnage par site d'étude.

Site	Superficie du site (km ²)	Nombre de stations	Densité (station/20 km ²)
R2001	5 625	274	0,97
R2002	1 401	81	1,16
S2002	1 340	76	1,13
T2002	1 372	73	1,06

Tableau 4 (suite)			
R2003-E	6 545	48	0,15
R2003-O	3 667	32	0,17
S2003	11 156	77	0,14
T2003	5 250	77	0,29
V2003	7 415	77	0,21

Méthodes de récolte des poils

La capture de poils d'ours noirs dans toutes ces études a été réalisée à l'aide de pièges à poils constitués de fils barbelés. Dans le cas de L2000, l'utilisation simultanée de pièges triangulaires et d'enclos a permis de tester l'efficacité de ces deux méthodes. Bien que les taux de visite pour les deux types d'installations aient été semblables, seul le piège triangulaire (figure 2) a été utilisé au cours des années suivantes, notamment à cause de sa rapidité d'installation et de son moindre coût (Lemieux et Desrosiers, 2002).

Les ours ont été attirés aux pièges à poils grâce à des gâteaux et des leurres olfactifs comme des essences de fruits, des chaudières de phoque en putréfaction ainsi que de l'huile de phoque (Courtois et coll., 2004). La récolte de poils s'est effectuée pendant l'été. Le tableau 5 présente les dates d'installation des stations et les dates des sessions de récolte des poils pour chaque étude. La récolte de poils s'est effectuée en cinq occasions en 2001 et en quatre occasions en 2002 et 2003. Les stations étaient visitées à des intervalles d'environ sept jours pour y collecter les poils et rafraîchir les appâts et leurres. Tous les poils présents sur le piège lors d'une visite étaient regroupés dans un même échantillon et conservés dans des enveloppes de papier. Les enveloppes contenant des poils ont été séchées à l'air libre et entreposées dans une boîte de carton aérée jusqu'aux analyses génétiques.

4.2. Analyses génétiques et démographiques

Analyses génétiques

Les analyses génétiques ont été réalisées par le laboratoire de génétique de l'Université du Québec à Rimouski (UQAR). Les méthodes d'analyse utilisées sont décrites dans Breton et Dufresne (2003). Les poils ont été examinés au microscope pour s'assurer qu'ils appartenaient à un ours et qu'ils possédaient une racine. Chaque échantillon analysé contenait au moins quatre poils et ceux contenant moins de quatre poils ou ne contenant aucun poil avec racine ont été écartés. L'extraction de l'ADN a été réalisée à l'aide du protocole Chelex (Breton et Dufresne, 2003). Quatre ou cinq marqueurs microsatellites ont été utilisés pour identifier les individus selon les années (tableau 6). L'efficacité de différents marqueurs microsatellites a été évaluée. Les échantillons d'ADN pour lesquels l'amplification sur gel révélait la présence de deux allèles différents pour le même locus étaient rejetés. Pour les génotypes n'apparaissant que dans un seul échantillon et ne différant que d'un seul locus, une deuxième lecture des gels a été réalisée pour confirmer que cet échantillon provenait bel et bien d'un nouvel individu et non d'une erreur de génotypage. Pour vérifier la constance des lectures, 15 échantillons tirés au hasard ont été analysés et relus (seulement pour l'année 2001). La probabilité de mauvaise identification n'est toutefois pas connue.

Tableau 5. Dates d'installation des stations d'échantillonnage, périodes et durées des sessions d'échantillonnage pour chaque étude.

Site d'étude	Période d'installation		Session 1	Session 2	Session 3	Session 4	Session 5
R2001	21 au 29 juil.	30 juil.	au 5 août	6 au 11 août	27 août au 2 sept.	3 au 9 sept.	10 au 15 sept.
R2002	Inconnue		5 au 11 août	12 au 18 août	19 au 24 août	25 au 30 août	
S2002	Inconnue		5 au 11 août	12 au 18 août	19 au 24 août	25 au 30 août	
T2002	Inconnue		5 au 11 août	12 au 18 août	19 au 24 août	25 au 30 août	
R2003-E et O	Inconnue		3 au 9 août	10 au 16 août	17 au 23 août	24 au 28 août	
S2003	Inconnue		3 au 9 août	10 au 16 août	17 au 23 août	24 au 28 août	
T2003	Inconnue		3 au 9 août	10 au 16 août	17 au 23 août	24 au 28 août	
V2003	Inconnue		3 au 9 août	10 au 16 août	17 au 23 août	24 au 28 août	

Tableau 6. Marqueurs microsatellites utilisés pour l'identification individuelle pour chaque inventaire.

	G1A	G1D	G10B	G10C	G10L	G10M	G10O	G10X
2001	X		X	X	X			
2002		X			X	X	X	
2003		X		X			X	X

Identification des génotypes identiques

Bien que nous ayons trouvé une liste des génotypes identifiés lors de chacun des inventaires, nous ne connaissons pas les méthodes utilisées pour les déterminer. Dans un souci de standardisation, nous avons donc utilisé la liste des génotypes issus des analyses génétiques pour refaire l'identification des génotypes identiques en utilisant le logiciel IDENTITY 4 (Wagner et Sefc, 1999). Ce logiciel apparie les génotypes dont les patrons alléliques sont identiques à partir d'une liste de génotypes complets sur un nombre de loci minimum désiré. Pour plus de détails, voir le mode d'emploi du logiciel IDENTITY 4 présenté à l'annexe 3. Dans le cadre du présent rapport, tous les génotypes complets sur trois loci ou plus ont été conservés et analysés.

Analyses démographiques

Les analyses démographiques visaient à estimer la taille et la densité des populations d'ours noirs selon les données recueillies dans le cadre des études décrites ci-haut. Nous nous attendions à obtenir des densités semblables à celles associées aux forêts résineuses et mélangées, soit entre 0,3 et 2,8 ours/10 km² respectivement (Boileau, 1993; Samson et Huot, 1994; Samson, 1996; Lamontagne et coll., 2006; Roy et coll., 2012). Nous avons testé deux approches différentes, soit 1) estimer la taille de population (N) et en déduire la densité en utilisant l'aire effective d'étude ($\hat{A}$), et 2) estimer directement la densité à l'aide de la maximisation de la vraisemblance pour des captures-recaptures spatialement explicites (ML SECR). Nous avons comparé l'efficacité de deux logiciels permettant de calculer la densité selon l'approche ML SECR, soit DENSITY 4.4 et R 2.12.2.

1) Estimation de la taille de population et de la densité

Pour estimer la taille de population (N), nous avons d'abord vérifié le degré de fermeture démographique des populations avec le logiciel CloseTest (Stanley et Richards, 2004). Pour plus de détails sur le fonctionnement de CloseTest, voir l'annexe 4. L'estimation de la taille des populations a été réalisée avec le logiciel MARK (White et Burnham, 1999).

Malgré l'étape précédente, nous avons tout de même privilégié les modèles de populations fermées, sans tenir compte des résultats du test de fermeture démographique. D'abord, il était raisonnable de présumer que les populations étaient plutôt fermées démographiquement et géographiquement pendant les études. En effet, les périodes d'échantillonnage étaient relativement courtes et à l'extérieur de la saison de chasse. L'utilisation de grandes aires d'étude relativement à la superficie moyenne des domaines vitaux et, dans certains cas, l'utilisation de barrières naturelles comme limites de la zone d'étude, ont permis de limiter l'effet de bordure et la violation de la supposition de fermeture géographique. L'utilisation de la superficie de l'aire d'étude effective ($\hat{A}$) permet de réduire l'importance de la supposition de fermeture géographique puisque toute la surface réellement couverte par les stations est considérée dans le calcul de la densité. Nous avons aussi opté pour les modèles de populations fermées puisque ces modèles permettent de considérer la variation individuelle, la réponse comportementale face à une première capture et l'effet du temps sur la probabilité de capture. Le logiciel MARK a été préféré au logiciel CAPTURE puisqu'il permet d'estimer la taille de population à l'aide d'une moyenne pondérée selon le poids relatif de chaque modèle. Cette méthode d'estimation était, selon nous, plus robuste que la sélection du modèle le plus performant à partir de l'AIC (Akaike, 1974) proposé dans CAPTURE. Nous avons choisi l'option des modèles complets de capture pour des populations fermées avec variation interindividuelle dans la probabilité de capture (*Full Closed Captures with Heterogeneity*). Nous n'avons pas utilisé l'option permettant d'intégrer les erreurs d'identification des génotypes dans l'estimation de l'abondance puisque celles-ci n'ont pas été évaluées et que, comme mentionné précédemment, les modèles intégrant la probabilité de mauvaise identification des individus ont de la difficulté à distinguer une mauvaise identification de la variation individuelle dans la probabilité de capture. La liste des modèles testés et leur signification biologique est présentée au tableau 2.

La taille de population a été estimée à l'aide du meilleur modèle lorsque ce dernier avait un poids supérieur à 90 % (Johnson et Omland, 2004). Sinon, l'estimation a été obtenue par le calcul de la moyenne pondérée selon le poids des modèles retenus (*model averaging*). Le mode d'emploi du logiciel MARK est résumé à l'annexe 5.

L'intervalle de confiance (IC) des estimations d'abondance issues de la moyenne pondérée a été calculé selon les équations suivantes :

$$IC = \left[M_{t+1} + (\hat{f}_0 / C), M_{t+1} + (\hat{f}_0 \times C) \right] \quad (\text{Équation 5})$$

$$\text{où} \quad \hat{f}_0 = \hat{N} - M_{t+1} \quad (\text{Équation 6})$$

$$\text{et} \quad C = \exp \left\{ 1.96 \left[\ln \left(1 + \frac{\text{var}(\hat{N})}{\hat{f}_0} \right) \right]^{1/2} \right\} \quad (\text{Équation 7})$$

où M_{t+1} correspond au nombre d'individus différents présents dans l'historique de capture, $\hat{N}$ est l'estimation d'abondance issue de la moyenne pondérée et $\text{var}(\hat{N})$ est la variance pouvant être obtenue en mettant au carré l'erreur type inconditionnelle fournie par MARK. Dans les cas où l'estimation d'abondance a été obtenue d'un seul modèle, nous avons utilisé la méthode traditionnelle de calcul de l'intervalle de confiance à l'aide de l'erreur type :

$$IC = \hat{N} \pm (1.96 \times \text{Erreur type}) \quad (\text{Équation 8})$$

La densité de population a été estimée grâce au calcul de l'aire d'étude effective ($\hat{A}$) dans le logiciel DENSITY 4.4 (Efford, 2009). Les quatre estimateurs d'aire effective disponibles (manuel, ARL/2, MMDM et MMDM/2) ont été utilisés et sommairement comparés entre eux. Pour l'estimateur nécessitant l'établissement manuel de la largeur de la bande entourant les stations (W), nous avons utilisé une valeur de 7 140 m qui correspond au diamètre d'un domaine vital circulaire de 20 km². Pour plus d'informations sur le fonctionnement du logiciel DENSITY 4.4, voir le mode d'emploi présenté à l'annexe 6. L'estimation de l'aire effective d'étude calculée avec ARL/2 a été préférée aux autres mesures lorsque disponible. Dans les cas où cette mesure ne pouvait être calculée dû à un nombre insuffisant de recaptures, nous avons utilisé l'estimation produite par la méthode manuelle. Nous avons préféré cette estimation à celles produites avec MMDM et MMDM/2 puisque ces dernières se basent aussi sur la détection des mouvements mis en évidence par

les recaptures entre les différentes stations. En plus des quatre estimations de l'aire effective, DENSITY offre aussi la possibilité de calculer l'aire effective à l'aide d'un polygone concave ou convexe entourant les stations d'échantillonnage (figure 4). L'option de calcul avec le polygone convexe est plus couramment utilisée, mais le polygone concave peut être plus approprié lorsque les stations sont positionnées de façon irrégulière et qu'elles sont éloignées les unes des autres (Efford, 2009).

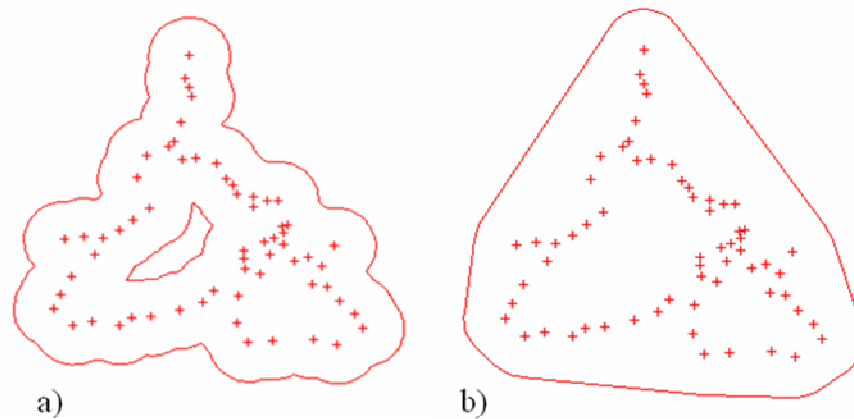


Figure 3. Représentation d'une aire effective d'étude utilisant a) un polygone concave et b) un polygone convexe.

2) Estimation de la densité avec la maximisation de la vraisemblance pour des captures-recaptures spatialement explicites

Pour la maximisation de la vraisemblance pour les captures-recaptures spatialement explicites (ML SECR), nous avons utilisé les logiciels DENSITY 4.4 et R 2.12.2. La liste des modèles testés et leur description sont présentées au tableau 3. Bien qu'il soit possible d'ajouter l'effet de la réponse comportementale et du temps dans la fonction dictant l'échelle spatiale à laquelle la probabilité de détection diminue (s), seule la variation individuelle dans la probabilité de capture est généralement modélisée (Drewry, 2010; Obbard et coll., 2010). La modélisation a été réalisée avec le type de stations « *proximity* » dans le logiciel DENSITY, qui permet à un individu d'être détecté à plus d'une station lors d'une session d'échantillonnage (Efford et coll., 2009), ce qui correspond bien à la méthode d'inventaire utilisée. Cependant, dans le logiciel R, le type de station « *proximity* » n'empêche pas que l'animal ait pu se déplacer suite à sa capture, mais il ne permet pas la

prise en charge de captures multiples d'un même individu dans la même session d'échantillonnage, comme cela a été le cas avec R2002 et S2002. Il faut alors utiliser le type de détecteur « *count* ».

Pour les deux logiciels, la maximisation de la vraisemblance des fonctions g et s est réalisée à partir de points discrets sur la grille d'échantillonnage. Dans le logiciel DENSITY, nous avons utilisé les valeurs établies par défaut, soit une matrice de 64×64 points, une distribution du nombre de détections par individu par station suivant une loi de Poisson uniforme et une distribution pour la fonction de détection (g) de forme demi-normale. Il est possible de réduire le nombre de points de cette matrice, mais cela peut compromettre la précision des résultats. Les estimations de densité de DENSITY sont accompagnées d'un intervalle de confiance asymétrique (pour plus de détails sur le calcul des intervalles de confiance, voir Borchers et coll., 2002). Pour plus d'informations sur le fonctionnement du logiciel DENSITY pour la réalisation de ML SECR, voir l'annexe 6.

Dans le logiciel R, nous avons utilisé l'extension *secl*. La distribution demi-normale pour la fonction de détection (g) a été utilisée et, par défaut, le logiciel utilise une distribution de Poisson pour la distribution du nombre de détections d'un individu à une station. Les modèles pour lesquels le logiciel suggérait que le biais dans l'estimation excédait 0,01 ont été abandonnés, de même que ceux montrant des messages d'erreurs suggérant que la maximisation de la vraisemblance avait échoué. Tout comme pour l'estimation de la taille de population, lorsque le meilleur modèle avait un poids AIC supérieur à 90 %, la densité était directement obtenue par ce dernier. Sinon, l'estimation de densité a été calculée à l'aide de la moyenne pondérée des modèles retenus. Pour plus d'informations sur le fonctionnement du logiciel R et de l'extension *secl*, voir l'annexe 7.

4.3. Résultats

Identification des génotypes semblables

Le nombre de paires de génotypes semblables, le nombre d'individus différents identifiés et le nombre résultant de recaptures sont présentés au tableau 7. Pour les sites R2003-O, S2003, T2003 et V2003, nous n'avons obtenu aucune paire d'individus semblables, donc aucun évènement de recapture, et pour les sites R2002, S2002 et T2002, nous avons obtenu un très faible nombre de recaptures. Ces sites ont donc été ignorés et nous avons poursuivi les analyses démographiques avec les données issues des inventaires de R2001 et R2003-E seulement. L'absence ou le très faible nombre de recaptures sont sans aucun doute dus à des erreurs lors des analyses génétiques. Le nombre d'individus différents dans le tableau 7 représente en fait le nombre d'échantillons pour lesquels l'analyse génétique a réussi, un nombre qui est souvent insuffisant pour révéler des évènements de recapture.

Tableau 7. Paires d'individus semblables, nombre d'individus différents identifiés et nombre de recaptures pour tous les sites d'étude. Les recaptures d'un même individu lors d'une même occasion de capture sont compilées.

Site d'étude	Paires semblables	Nombre d'individus différents	Nombre de recapture
R2001	143	483	143
R2002	10	124	8
S2002	6	75	6
T2002	5	147	5
R2003-O	0	21	0
R2003-E	21	46	16
S2003	0	32	0
T2003	0	38	0
V2003	0	34	0

Fermeture démographique des populations

Le test de fermeture démographique de population a pu être réalisé sur une seule base de données, soit celle de R2001. Le logiciel a établi que la population était ouverte. Cela suggère qu'il faut un nombre assez élevé de captures et de recaptures pour réaliser ce test.

Estimation de l'abondance

Plusieurs problèmes ont été rencontrés lors de l'estimation de la taille des populations dans le logiciel MARK. D'abord, certains modèles généraient le message d'avertissement suivant : *Warning ** error number 2 from VA09AD optimization routine*. Ce message indique qu'un problème est survenu lors de l'optimisation des valeurs des paramètres du modèle. Bien que ce message d'erreur ne soit apparu que pour quelques modèles, il est possible que la routine d'optimisation des valeurs des paramètres se soit terminée sans message d'erreur pour les autres modèles, mais que les résultats ne soient pas fiables. Aussi, certains modèles étaient vraisemblablement surparamétrisés, donnant ainsi des résultats d'abondance anormalement élevés. En fait, les bases de données ne contenaient parfois pas suffisamment de données, particulièrement de recaptures, pour estimer tous les paramètres des modèles avec variation individuelle de la probabilité de capture (Gary White, Université du Colorado, commun. pers.). Les résultats des analyses et les problèmes rencontrés sont présentés dans les sections suivantes.

a) Résultats pour le site de Rouyn-Noranda en 2001 (R2001)

Les modèles incluant une influence du comportement sur la probabilité de capture se sont avérés problématiques. Nous avons donc retiré les modèles incluant la réponse comportementale b de la liste de modèles candidats, et nous avons recalculé les poids des modèles restants (tableau 8). Puisque le poids du modèle M_{th} était de 0,97, nous avons utilisé l'estimation issue de ce modèle pour estimer la densité (1 208 ours [902; 1 517]).

Tableau 8. Estimation de la taille de population pour R2001 selon les modèles de populations fermées dans MARK. Les estimations et les erreurs types en caractère gras sont celles issues des modèles pour lesquels l'estimation a vraisemblablement échoué.

Modèle	Estimation N	Erreur type	AICc	$\Delta AICc$	Poids du modèle	Nombre de paramètres	Déviance
M_{tbh}	3137407,3	0,0	-2859,0	0,0	0,76	9	56,27
M_{tb}	465455,6	5246528,2	-2855,4	3,5	0,13	7	63,83
M _{th}	1208,9	156,8	-2855,1	3,9	0,11	8	62,19
M _t	1056,8	75,7	-2847,7	11,2	0,00	6	73,53
M_{bh}	1937776,2	12497835,0	-2802,9	56,1	0,00	5	120,44
M_b	1111386,1	0,0	-2795,4	63,6	0,00	3	131,93
M _h	1242,9	165,8	-2728,2	130,7	0,00	4	197,06
M _o	1099,1	80,2	-2722,8	136,2	0,00	2	206,54

b) Résultats pour le site de Rouyn-Noranda Est en 2003 (R2003-E)

Le paramètre de réponse comportementale semble aussi avoir fait échouer l'estimation de la taille de population. Alors que les modèles M_{bh} et M_b ont produit des estimations d'abondance et des erreurs types irréalistes, les modèles M_{tb} et M_{tbh} ont fourni des estimations plus réalistes mais pour lesquelles les erreurs types étaient plus élevées que les estimations elles-mêmes. Nous avons donc retiré tous les modèles contenant la variable *b* de la liste des modèles candidats pour le calcul de la moyenne pondérée (tableau 9). Bien que les modèles M_h et M_{th} aient donné des estimations réalistes, ils ont produit des estimations et des erreurs types identiques aux modèles M_o et M_t, respectivement. Ceci suggère que la variable *h* n'a pas été prise en compte pour l'estimation de la taille de population. Les modèles contenant la variable *h* ont donc aussi été mis de côté pour le calcul de la moyenne pondérée. La taille de population pour le site de R2003-E, calculée à partir des modèles M_o (0,58) et M_t (0,16), était de 113 individus ([74; 208]).

Tableau 9. Estimation de la taille de population pour R2003-E selon les modèles de populations fermées dans MARK. Les estimations et les erreurs types en caractère gras sont celles issues des modèles pour lesquels l'estimation de l'abondance a vraisemblablement échoué.

Modèle	Estimation N	Erreur type	AICc	Δ AICc	Poids du modèle	Nombre de paramètres	Déviance
M_{bh}	118745,7	4156428,9	-76,7	0,0	0,49	2	14,74
M _o	113,6	30,6	-74,7	2,0	0,18	2	16,69
M_b	768656,6	14767928,0	-74,6	2,1	0,17	3	14,74
M _h	113,6	30,6	-72,6	4,0	0,07	3	16,69
M _t	111,6	29,8	-72,1	4,6	0,05	5	13,04
M_{tb}	91,3	114,0	-70,0	6,7	0,02	6	13,02
M _{th}	111,6	29,8	-70,0	6,7	0,02	6	13,04
M_{tbh}	91,3	114,0	-67,8	8,9	0,01	7	13,02

Estimation de la superficie effective du site d'étude

Les superficies effectives des aires d'étude calculées avec l'estimateur manuel et l'ARL/2, ainsi que les deux types de polygones (convexe ou concave), sont présentées dans le tableau 10.

Estimation de la densité de population à partir de l'abondance et de l'aire effective d'étude

Nous avons utilisé une bande entourant les stations d'une largeur de 5 046 m, soit le diamètre d'un domaine vital de 20 km². La densité a été estimée en divisant l'abondance N obtenue dans DENSITY par l'aire d'étude effective $\hat{A}$ issue du calcul avec les polygones concave et convexe (tableau 11). La densité estimée à partir du polygone concave était systématiquement plus élevée que celle issue du calcul avec un polygone convexe. Cependant, la différence entre les estimations obtenues par le polygone concave et le polygone convexe n'était pas significative pour R2001.

Estimation de la densité avec la maximisation de la vraisemblance pour des captures-recaptures spatialement explicites (ML SECR) dans le logiciel DENSITY 4.4

Les estimations de densité obtenues dans le logiciel DENSITY 4.4 sont aussi variables, mais dans l'ensemble plus réalistes que celles dérivées des estimations d'abondance et des

aires effectives. Pour certains sites, les estimations de densité obtenues sont anormalement élevées, suggérant des problèmes à déterminer la valeur de certains paramètres des modèles. Les résultats de densité et leurs problèmes associés sont présentés selon l'année d'étude.

a) Site Rouyn-Noranda en 2001 (R2001)

Il a été impossible d'obtenir les résultats des analyses dans DENSITY 4.4 pour R2001 après plusieurs jours d'attente, possiblement à cause de la complexité des calculs. Les analyses ont cependant fonctionné dans l'environnement de travail R.

Tableau 10. Superficie effective de l'aire d'étude (km²) pour tous les sites, selon l'estimateur manuel (bande de 5 056 m), l'ARL/2 et les polygones convexe et concave. Les nombres entre parenthèses correspondent à la largeur (m) de la bande entourant les stations d'échantillonnage.

Manuelle (5 046 m)		ARL/2		MMDM/2		MMDM	
Site	Polygone convexe	Polygone concave	Polygone convexe	Polygone concave	Polygone convexe	Polygone concave	Polygone concave
R2001	7 126,08	6 362,60	8 460,98 (9 140)	7 901,26 (9 140)	5 796,84 (604)	162,40 (604)	5 970,32 (1 208)
R2003-E	8 281,66	1 625,80	7 549,85 (2 976)	777,36 (2 976)	2 109,49 (4 691)	1 777,59 (4 691)	2 984,88 (9 381)

Tableau 11. Densité d'ours (individu/10 km²) estimée dans chaque site d'étude avec l'approche N/Â à partir de l'aire effective d'étude calculée en utilisant la méthode du polygone concave ou du polygone convexe, et proportion des individus capturés une seule fois selon les génotypes.

Site d'étude	Polygone concave		Polygone convexe		Proportion des individus capturés une seule fois
	Densité (ind./10 km ²)	Intervalle de confiance	Densité (ind./10 km ²)	Intervalle de confiance	
R2001	1,53	[-8,15; 11,20]	1,43	[-7,61, 10,46]	76,90
R2003-E	1,46	[0,51; 4,13]	0,15	[0,05, 0,43]	74,19

* Densité calculée avec une aire d'étude effective établie en utilisant une bande de 5 046 m autour des stations.

b) Site Rouyn-Noranda en 2003 (R2003-E)

Pour R2003-E, tous les modèles ayant un poids supérieur à 0 ont été conservés pour le calcul de la moyenne pondérée de la densité (tableau 12). L'estimation de la densité est de 1,59 ours/10 km² ([0,81; 1104,39]). Cette estimation est réaliste et concorde avec la densité attendue pour cette région. Toutefois, l'intervalle de confiance pondéré est anormalement large dû à l'estimation issue du modèle $g(h)s(.)$ qui semble incohérente. Ceci suggère, pour une raison inconnue, un échec du logiciel à modéliser correctement la densité avec ce modèle. Nous avons donc retiré celui-ci de la liste pour le calcul de la densité pondérée, et l'estimation de densité ainsi obtenue est de 1,50 ours/10 km² ([0,81; 2,76]).

Tableau 12. Densité d'ours estimée (individu/10 km²) pour le secteur R2003-E pour chacun des modèles avec DENSITY.

Modèle	AICc	$\Delta AICc$	Poids du modèle	Densité	Erreur type	Intervalle de confiance	Densité pondérée
$g(t)s(.)$	358,5	0	0,35	1,69	0,52	[0,94; 3,04]	0,60
$g(tb)s(.)$	358,9	0,4	0,28	1,18	0,40	[0,62; 2,24]	0,33
$g(.)s(.)$	360,1	1,6	0,16	1,70	0,52	[0,94; 3,06]	0,26
$g(b)s(.)$	362,5	4,0	0,05	1,66	0,63	[0,81; 3,39]	0,08
$g(th)s(.)$	362,9	4,4	0,04	1,69	0,52	[0,94; 3,04]	0,07
$g(t)s(h)$	363,1	4,6	0,03	1,69	0,51	[0,94; 3,00]	0,06
$g(tb)s(h)$	364,3	5,8	0,02	1,17	0,39	[0,62; 2,20]	0,02
$g(tbh)s(.)$	364,3	5,9	0,02	1,18	0,39	[0,6; 0,22]	0,02
$g(.)s(h)$	365,1	6,6	0,01	1,74	0,53	[0,97; 3,13]	0,02
$g(th)s(h)$	366,4	7,9	0,01	1,68	0,51	[0,94; 3,00]	0,01
$g(tbh)s(h)$	366,7	8,3	0,01	1,19	0,40	[0,62; 2,27]	0,01
Moyenne et intervalle de confiance pondérés							1,50 [0,81; 2,76]

Estimation de la densité avec la maximisation de la vraisemblance pour des captures-recaptures spatialement explicites (ML SECR) dans le logiciel R.

La maximisation de la vraisemblance a échoué pour certains modèles lors de l'estimation de la densité avec l'approche ML SECR dans le logiciel R. L'échec des ML SECR s'est avéré plus fréquent pour les modèles avec un nombre élevé de paramètres à évaluer, suggérant que le nombre de captures et de recaptures était trop faible pour estimer correctement la densité avec cette méthode.

a) Site Rouyn-Noranda en 2001 (R2001)

Contrairement au logiciel DENSITY, le logiciel R est parvenu à réaliser avec succès la ML SECR pour R2001. Toutefois, l'exécution du modèle a accaparé un ordinateur personnel pendant près de deux semaines. D'ailleurs, les deux modèles ayant le plus grand nombre de paramètres, $g(tbh)s(.)$ et $g(tbh)s(h)$, ont dû être abandonnés. Le tableau 13 présente les résultats de l'estimation de la densité avec l'approche ML SECR dans R pour R2001. Les estimations de densité étaient plus élevées que ce qui est attendu pour la forêt résineuse, en plus d'être peu précises. La densité estimée par la moyenne pondérée est de 8,12 ours/10 km² ([4,73; 14,07]).

Tableau 13. Densité d'ours estimée (individu/10 km²) pour le secteur R2001 pour chacun des modèles testés avec l'approche ML SECR dans R.

Modèle	AICc	$\Delta AICc$	Poids du modèle	Densité	Erreur type	Intervalle de confiance	Densité pondérée
$g(tb)s(h)$	2 950,3	0,0	0,84	8,62	2,59	[4,84; 15,34]	7,28
$g(t)s(h)$	2 953,7	3,4	0,16	5,43	0,77	[4,12; 7,15]	0,84
$g(h)s(h)$	3 072,1	121,8	0,00	4,92	7,03	[3,72; 6,50]	0,00
$g(.)s(h)$	3 083,7	133,4	0,00	5,55	0,81	[4,18; 7,38]	0,00
$g(th)s(.)$	3 268,3	318,0	0,00	2,63	0,35	[2,03; 3,41]	0,00
$g(t)s(.)$	3 305,6	355,3	0,00	1,96	1,49	[1,69; 2,27]	0,00
$g(bh)s(.)$	3 338,6	388,3	0,00	10,51	3,68	[5,40; 20,46]	0,00
$g(b)s(.)$	3 366,8	416,5	0,00	7,44	2,21	[4,21; 13,16]	0,00
$g(h)s(.)$	3 402,5	452,2	0,00	2,61	0,34	[2,02; 3,37]	0,00
$g(.)s(.)$	3 428,0	477,6	0,00	1,98	0,15	[1,70; 2,30]	0,00
Moyenne et intervalle de confiance pondérés							8,12 [4,73; 14,07]

b) Site Rouyn-Noranda Est en 2003 (R2003-E)

Pour R2003-E, plusieurs modèles ont échoué à estimer la densité, bien que le logiciel ait suggéré une bande limitrophe aux stations de largeur raisonnable (3 969 m). Le tableau 14 présente les estimations de densité pour le site R2003-E obtenues par ML SECR dans le logiciel R. Puisque le modèle le plus performant avait un poids inférieur à 90 %, nous avons effectué le calcul de la moyenne et de l'intervalle de confiance pondérés avec tous les

modèles ne présentant pas de problème apparent. La densité pondérée obtenue est de 3,14 individus/10 km² ([1,54; 6,47]), ce qui est réaliste pour le type d'habitat échantillonné.

Tableau 14. Densité d'ours estimée (individu/10 km²) pour le secteur R2003-E pour chacun des modèles avec l'approche ML SECR dans R.

Modèle	AICc	$\Delta AICc$	Poids du modèle	Densité	Erreur type	Intervalle de confiance	Densité pondérée
g(.)s(.)	359,7	0,0	0,30	1,82	0,57	[1,01; 3,30]	0,55
g(.)s(h)	359,8	0,1	0,29	4,96	1,97	[2,35; 10,50]	1,42
g(t)s(.)	361,8	2,1	0,10	1,81	0,56	[1,00; 3,27]	0,19
g(b)s(.)	362,1	2,4	0,09	1,76	0,65	[0,88; 3,55]	0,16
g(b)s(h)	362,4	2,7	0,08	4,86	2,28	[2,03; 11,63]	0,37
g(t)s(h)	362,4	2,8	0,08	4,93	1,91	[2,37; 10,27]	0,37
g(tb)s(.)	362,6	2,9	0,07	1,30	0,43	[0,69; 2,44]	0,09
Moyenne et intervalle de confiance pondérés							<u>3,14 [1,54; 6,47]</u>

Comparaison des estimations de densité obtenues avec les trois approches utilisées : aire effective, ML SECR dans le logiciel DENSITY 4.4 et ML SECR dans le logiciel R

Dans une démarche exploratoire, nous avons comparé les estimations de densité obtenues à partir de l'estimation d'abondance et de l'aire d'étude effective (c.-à-d. $N/\hat{A}$ avec le polygone concave) et celles issues de la maximisation de la vraisemblance pour des captures-recaptures spatialement explicites. Nous avons utilisé un simple test de Z pour vérifier si les estimations de densité étaient équivalentes. Nous avons ensuite comparé la précision des estimations de densité. Nous avons calculé la précision (P) comme suit :

$$P = \frac{\left(\frac{(IC95\% \lim_{sup} - IC95\% \lim_{inf})}{D} \times 100 \right)}{2} \quad (\text{Équation 9})$$

où D représente l'estimation de densité, et $\lim_{sup}$ et $\lim_{inf}$ représentent respectivement les limites supérieures et inférieures de l'intervalle de confiance.

Le tableau 15 présente les estimations de densité, les erreurs types, les intervalles de confiance ainsi que les coefficients de variation modifiés pour tous les sites d'étude selon les deux approches ($N/\hat{A}$ et ML SECR). Dans le cas de l'approche ML SECR, les résultats présentés sont ceux obtenus par les deux logiciels utilisés, à savoir DENSITY et R.

a) Site Rouyn en 2001 (R2001)

Seules les estimations provenant des $N/\hat{A}$ et ML SECR dans R ont été comparées puisque le calcul de la ML SECR n'a pas pu être réalisé dans DENSITY pour ce site. L'estimation issue de la ML SECR dans R était significativement plus élevée que celle obtenue avec $N/\hat{A}$. La précision des estimations était mauvaise, et inférieure à celle recherchée (25 %), mais elle était près de 11 fois meilleure avec l'approche ML SECR dans R (57,5 %) qu'avec l'approche $N/\hat{A}$.

b) Site Rouyn Est en 2003 (R2003-E)

Pour ce site, les trois estimations de densité obtenues par $N/\hat{A}$ et ML SECR dans le logiciel DENSITY et par ML SECR dans R étaient statistiquement équivalentes. Cependant, l'estimation de densité obtenue par ML SECR dans R est celle dont la précision était la meilleure (78,5%), bien qu'inférieure à la précision recherchée.

Comparaison des estimations d'abondance et de densité obtenues dans cette étude avec celles de Courtois et coll. (2004) pour le site de R2001

Pour estimer la taille de population à partir des données de R2001, Courtois et coll. (2004) ont utilisé les 10 estimateurs (modèles) disponibles dans le programme CAPTURE auxquels ils ont ajouté les estimateurs de Peterson et Chapman. Le test d'ajustement (*goodness-of-fit*) a été employé pour déterminer quel estimateur s'ajustait le mieux aux données. Ils ont jugé pour ce jeu de données que l'effet de bordure était négligeable vu la taille importante du bloc d'étude (estimé à 5 550 km²) par rapport à la superficie des domaines vitaux des femelles (20 km²).

Tableau 15. Comparaison des estimations de densité calculées à partir des estimations d'abondance et de la superficie de l'aire effective d'étude ($N/\hat{A}$) et celles issues de la maximisation de la vraisemblance de captures-recaptures spatialement explicites (ML SECR) dans DENSITY et dans R. Les erreurs types (S.E.), les intervalles de confiance (IC95 %) et la précision (%) sont aussi présentés pour chaque site d'étude. Les densités portant un astérisque ont été calculées avec une bande entourant les stations de 5 046 m de largeur (établie manuellement).

	N/Â			ML SECR dans DENSITY				ML SECR dans R				
	Densité	S.E.	IC95 %	Précision (%)	Densité	S.E.	IC95 %	Précision (%)	Densité	S.E.	IC9 5%	Précision (%)
Site												
R2001	1,53	0,19	[-8,15; 11,20]	632,4	n/d	n/d	n/d	n/d	8,12	2,31	[4,73; 14,07]	57,5
R2003-E	1,46	0,39	[0,51; 4,13]	124,0	1,50	0,4776	[0,81; 27,56]	891,7	3,14	1,19	[1,54; 6,47]	78,5

Tableau 16. Tableau tiré de Courtois et coll. (2004) présentant l'estimation de l'abondance, les erreurs types (S.E.) et les intervalles de confiance (IC95 %) pour les historiques de capture surestimant le nombre probable de recaptures, estimant le plus correctement possible le nombre probable de recaptures et sous-estimant le nombre probable de recaptures (pour plus de détails sur les différents critères de sélection des événements probables de recapture, voir Courtois et coll. [2004]).

Fichier utilisé	n génotypes	Modèle	Nombre d'ours estimé		
			Population	S.E.	IC95 %
Recaptures surestimées	312	M_{th}	468	33	416 - 545
Recaptures probables	392	M_h	817	40	745 - 903
Recaptures sous-estimées	450	M_{th}	1184	128	973 - 1481

Ils ont aussi considéré que les oursons n'étaient pas échantillonnés avec le type de piège utilisé puisque ces derniers étaient trop petits pour atteindre l'appât et que la mère était susceptible de monopoliser l'appât. Le modèle retenu par le logiciel CAPTURE était M_{th} (Chao) et l'abondance pour la population totale était estimée à 1 103 ours ([995, 1211]) [tableau 19]. Rappelons que l'estimation d'abondance que nous avons obtenue dans MARK était de 1 208 ours ([901, 1516]). Ces estimations peuvent être considérées comme équivalentes, bien que l'estimation de Courtois et coll. (2004) soit plus précise. Toutefois, dans le cadre de leur rapport, Courtois et coll. (2004) ont considéré que la méthode d'attribution des génotypes identiques la plus réaliste, dont dépend le nombre de recaptures, était celle éliminant toutes les recaptures étant survenues à plus de 15 km du lieu de la capture, ce qui a mené à une estimation finale de 817 ours ([745, 903]) [tableau 16]. La densité dérivée de cette estimation d'abondance est de 1,50 ours/10 km² ([1,4; 1,6]), ce qui est presque égal à celle obtenue par notre approche $N/\hat{A}$ de 1,53 ours/10 km² ([1,2; 2,0]), mais très différent de la densité obtenue par la méthode ML SECR dans R (8,12 individus/10 km², [4,7; 14,1]).

5. ESTIMATION DE LA DENSITÉ SELON DIFFÉRENTS SCÉNARIOS D'ÉCHANTILLONNAGE PAR SIMULATION

5.1. Objectifs et méthodologie des simulations

Un inventaire de l'ours noir utilisant la reconnaissance individuelle par analyse génétique des poils implique un investissement important en ressources humaines et financières. Par exemple, le succès d'un tel projet demande d'obtenir un nombre assez élevé de captures et de recaptures. Or, le nombre de recaptures augmente avec le nombre de stations d'échantillonnage, mais plus de stations implique une période d'installation plus longue, plus d'échantillons de poils, des analyses génétiques plus dispendieuses, etc. Il est par conséquent important de trouver un compromis acceptable entre le coût de l'opération et la précision des résultats obtenus. Nous avons donc établi la relation entre le nombre de pièges et la précision des estimations de densité à partir de la base de données de R2001. Il

n'a pas été possible de faire le même exercice avec les données de R2003-E, car il n'y avait pas assez de stations disponibles ($n = 32$).

En utilisant cette base de données, nous avons estimé la densité pour différents scénarios en faisant varier le nombre de stations et nous avons comparé la précision des estimations issues de ces simulations. Nous avons retiré aléatoirement de 10 à 90 % (par tranche de 10 %) des stations parmi celles disponibles, en créant une nouvelle base de données à chaque itération. Nous avons ajusté l'historique de capture en fonction de l'identité des stations retirées. Le tableau 17 présente le nombre de captures, de recaptures et de stations conservées dans les simulations. Nous avons calculé la densité en utilisant l'approche ML SECR dans le logiciel statistique R. La largeur de la bande W utilisée était celle suggérée par le logiciel R (fonction *suggest.buffer*). Les modèles incluant la variable de réponse comportementale à une première capture b n'ont pas été utilisés vu la piètre qualité des résultats qu'ils ont fournis pour presque toutes les analyses effectuées avec MARK, DENSITY et R. Les modèles pour lesquels le logiciel R suggérait que le biais excédait 1 % de la densité prédite ont été abandonnés, de même que ceux montrant des messages d'erreur indiquant l'échec de la maximisation de la vraisemblance. La densité pour chaque base de données a été calculée avec la moyenne pondérée de tous les modèles lorsque le meilleur modèle avait un poids inférieur à 0,90. Dans le cas contraire, seul le meilleur modèle a été utilisé pour l'estimation.

Tableau 17. Nombre de captures, de recaptures et de stations conservées pour chacune des bases de données utilisées dans les simulations permettant d'estimer la densité d'ours à partir d'un sous-échantillon des stations.

Pourcentage de stations conservées	Nombre de stations	Nombre de captures	Nombre d'individus	Nombre de recaptures
10	27	59	48	9
20	55	128	91	37
30	81	186	149	37
40	110	253	190	63
50	136	319	255	64
60	164	366	296	70
70	192	444	347	97
80	219	496	387	109
90	247	553	428	125

Comme suggéré par Garshelis (2011), les estimations de population ou de densité servant à la gestion des populations devraient avoir une probabilité de 95 % ou se trouver à ± 25 % de la taille ou de la densité réelle de la population. Autrement dit, les limites inférieures et supérieures des intervalles de confiance (à un seuil $\alpha = 95$ %) ne devraient pas être inférieures ou supérieures à 25 % de l'estimation. Nous avons donc évalué la précision des estimations en calculant le pourcentage de déviation des limites inférieures et supérieures des intervalles de confiance par rapport à l'estimation. Nous avons utilisé une régression linéaire pour évaluer la relation entre la proportion de stations conservées et la déviation des limites inférieures et supérieures des estimations de densité.

5.2. Résultats des simulations

La maximisation de la vraisemblance a échoué pour les bases de données pour lesquelles seulement 10 % et 20 % des stations ont été conservées, probablement à cause d'un nombre insuffisant d'évènements de capture et de recapture (tableau 18). En effet, le nombre de recaptures est identique pour les bases de données avec 20 % et 30 % des stations, toutefois, le nombre d'individus différents identifiés est largement supérieur pour celles avec 30 % des stations.

Tableau 18. Estimation de la densité, erreur type et identification du meilleur modèle selon la proportion de stations conservées dans la base de données de R2001.

Pourcentage de stations conservées (%)	Densité	Erreur type	Meilleur modèle
30	6,41	5,68	g(th)s(h)
40	4,76*	1,33*	n/d
50	8,61	2,40	g(t)s(h)
60	4,82	0,81	g(th)s(h)
70	5,06*	0,91*	n/d
80	7,57*	1,47*	n/d
90	5,27	0,76	g(t)s(h)
100	5,43	0,77	g(t)s(h)

* Estimations issues du calcul de la moyenne pondérée des modèles candidats.

La précision des estimations de densité diminuait lorsque le nombre de stations conservées augmentait (figure 5). L'équation 10 décrit la relation linéaire négative entre la proportion de stations conservées (x) et la déviation des limites inférieures des intervalles de confiance (y) :

$$y = 326,14 - 6,46x + 0,04x^2 \quad (\text{Équation 10})$$

En posant $y = 50\%$ et en résolvant l'équation 10, il est possible de déterminer le nombre minimum de stations requises pour obtenir une erreur maximale de 25 % de part et d'autre de l'estimation, c'est-à-dire que les limites inférieure et supérieure de l'intervalle de confiance ne sont pas à plus de 25 % de l'estimation de densité. La résolution mène à l'obtention de $x = 81,1\%$ des stations originalement utilisées pour R2001, soit 222 stations. Ainsi, un minimum de 222 stations serait requis lors de l'inventaire de R2001 pour maximiser les chances d'obtenir, en moyenne, le niveau de précision recherché pour la gestion des populations.

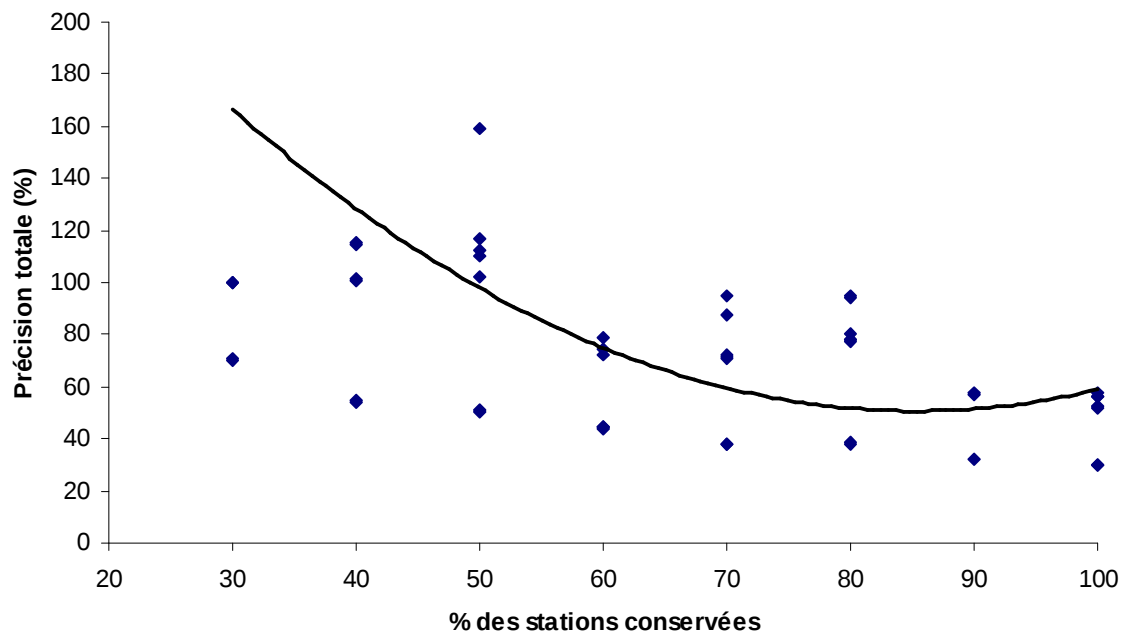


Figure 4. Relation entre la précision de l'estimation de densité (étendue de l'intervalle de confiance divisée par l'estimation) et le pourcentage (%) de stations conservées pour l'estimation de la densité pour R2001.

6. DISCUSSION ET RECOMMANDATIONS

6.1 Nombre et disposition des stations d'échantillonnage

Les simulations ont démontré qu'un minimum de 222 stations d'inventaire aurait permis d'obtenir, en moyenne, une précision acceptable pour la gestion des populations dans le cadre de l'inventaire de R2001. Cependant, un examen plus attentif des données issues des simulations indique que, même avec un tel échantillon, plusieurs estimations n'avaient pas la précision recherchée. Dans ce contexte, nous croyons que 250 stations constitueraient un minimum plus conservateur. Ce nombre est supérieur au nombre de stations généralement utilisées dans les études semblables. La disposition des stations d'échantillonnage est un élément important dans la réussite de ce type d'inventaire. Dans le cadre des inventaires rapportés dans le présent rapport, les stations étaient positionnées de façon assez uniforme selon une densité de 1,0 à 1,2 station/20 km². Le site de R2003-E est le seul site pour lequel la disposition des stations était en grappes avec une densité de 0,15 station/20 km² (les stations étaient aussi distribuées en grappes pour S2003, T2003 et V2003, mais la densité n'a pas pu être estimée pour ces sites). Malgré la faible densité des stations pour R2003-E, le taux de recapture était le deuxième plus élevé parmi les sites, ce qui pourrait s'expliquer par la disposition en grappes des stations. Cette observation concorde avec Roy et coll. (2007) qui ont rapporté que l'utilisation de cinq stations par parcelle plutôt qu'une seule a permis de capturer de nouveaux individus, particulièrement des femelles. Cette approche permet de maximiser les chances d'avoir des recaptures et de diminuer les coûts reliés au temps requis pour les visites hebdomadaires. Nous croyons que la disposition en grappes pourrait être avantageuse, particulièrement pour couvrir un grand territoire. Cependant, une disposition en grappes des stations ne permet pas de couvrir une superficie totale aussi grande qu'une disposition uniforme. De plus, la prudence est de mise lorsqu'on doit extrapoler la densité à une grande aire d'étude alors que seulement une faible proportion de celle-ci a été échantillonnée (Boulanger et coll., 2004).

6.2 Efficacité de la méthode de capture et de récolte des poils

Le piège triangulaire a été préféré à l'enclos de fils barbelés dans les études réalisées en Abitibi-Témiscamingue (Courtois et coll., 2004; Roy et coll., 2007). Cette méthode a été retenue, car elle est moins dispendieuse et les stations sont plus rapides à installer (Lemieux et Desrosiers, 2002) et elle a aussi fait ses preuves ailleurs en Amérique du Nord (p. ex., Woods et coll., 1999; Boersen et coll., 2003). Nous croyons cependant que le piège triangulaire a le désavantage d'offrir une faible superficie de barbelés aux ours qui se présentent à l'appât. Ceci pourrait augmenter considérablement le risque de contamination des échantillons (ours multiples), notamment dans les régions où la densité d'ours est élevée et où les chances sont élevées d'avoir la visite de plusieurs ours dans la même semaine. Nous recommandons donc l'utilisation d'enclos de barbelés (environ 30 m de circonférence) dans les cas où la contamination des échantillons de poils (ours multiples) est probable.

Quant aux méthodes de récolte des poils, les rapports existants sur les inventaires faisant l'objet du présent rapport ne font pas mention de leur efficacité. Toutes ces études ont utilisé la dessiccation dans une enveloppe de papier. Il est important de porter des gants lors de la récolte des poils pour éviter toute contamination de l'ADN, particulièrement si les marqueurs microsatellites utilisés pour le génotypage sont communs à l'ours noir et à l'humain ou si le sexage des individus est réalisé par un autre protocole que celui de Yamamoto et coll. (2002).

Il est difficile de dire si le regroupement des échantillons de poils retrouvés sur un même piège lors d'une visite hebdomadaire est une bonne idée en se basant sur les résultats obtenus dans le présent rapport. Toutefois, Tredick et coll. (2007) ont montré que les échantillons de poils retrouvés sur des pointes de barbelés séparées par une pointe d'intervalle provenaient d'individus différents dans 23 à 55 % des cas. Ainsi, nous suggérons que des échantillons de poils séparés par une pointe de fil barbelé « vide » soient considérés comme des échantillons différents. Cette méthode implique cependant une

augmentation du nombre total d'échantillons de poils qui devront être analysés et, conséquemment, des coûts reliés aux analyses génétiques.

6.3 Efficacité des appâts et des leurres olfactifs

Les retailles de gâteaux commerciaux combinées à un vaporisateur commercial ou composé d'huile de phoque et de chair de phoque en décomposition ont permis d'attirer les ours noirs aux stations. Ces leurres et appâts sont généralement abordables et accessibles (voir Lemieux et Desrosiers [2002], pour une liste exhaustive du matériel utilisé et des coûts pour l'étude de L2000). Cependant, afin de respecter les suppositions de l'approche de CMR, il faut éviter d'utiliser un appât qui serait susceptible de retenir les ours à proximité des stations. À cet effet, l'utilisation d'un simple leurre olfactif (p. ex., huile de poisson, huile essentielle d'anis, sang, etc.) devrait être privilégiée. Si un appât est utilisé, il devrait être en quantité suffisamment petite pour ne pas constituer une récompense intéressante pour l'ours.

6.4 Période et durée de l'échantillonnage

Les périodes d'échantillonnage des études décrites dans le présent rapport se sont étendues du début juillet à la mi-septembre (tableau 5). Certains auteurs ont suggéré de réaliser les inventaires plus tôt dans la saison estivale pour éviter d'échantillonner les longues excursions de quête alimentaire qui se produisent généralement à la fin de l'été (prémisse de fermeture géographique violée; Rogers, 1987) et de profiter de la rareté des petits fruits en juin pour maximiser le taux de capture aux stations (Gignac et coll., 2008). Il est donc possible que les inventaires analysés dans le présent rapport aient enregistré des mouvements de grande envergure à cause de la période assez tardive à laquelle les travaux ont été réalisés. Nous suggérons que l'inventaire des populations d'ours noirs soit réalisé tout de suite après la saison de chasse du printemps (en juin ou en juillet) afin de minimiser les chances que les animaux effectuent de grands déplacements reliés à la quête alimentaire, ce qui pourrait compromettre le respect de la supposition de fermeture géographique de la population. De plus, concentrer l'effort d'échantillonnage sur quatre semaines pourrait

limiter les chances de violation des suppositions de fermeture démographique et géographique.

6.5 Analyses génétiques

Le succès des analyses génétiques est essentiel à la réussite de ce type d'inventaire. Les études décrites dans le présent rapport ont utilisé quatre ou cinq marqueurs microsatellites. Le choix du nombre de microsatellites repose sur la probabilité d'identité, qui varie grandement dans la littérature, par exemple de $6,08 \times 10^{-14}$ avec 13 marqueurs microsatellites (Robinson et coll., 2007) à 0,831 avec 8 marqueurs (Boersen et coll., 2003). Nous croyons qu'il serait avantageux de faire une étude préliminaire visant à établir le nombre et l'identité des marqueurs optimaux pour maximiser les chances de différenciation individuelle. Une façon d'aborder le problème du nombre de marqueurs microsatellites à utiliser est de considérer la méthode d'analyse et surtout de réanalyse des génotypes semblables. Paetkau (2004) recommande de réanalyser tous les génotypes qui diffèrent par trois microsatellites ou moins (nommés 3MM-pairs). Cette méthode est plus conservatrice par rapport aux méthodes de réanalyse précédentes (1MM-pairs et 2MM-pairs; Woods et coll., 1999; Paetkau, 2003), et elle permet d'augmenter la précision des appariements de génotypes tout en utilisant un nombre limité de marqueurs microsatellites. En s'appuyant sur les recommandations de Paetkau (2004), l'utilisation de cinq à sept marqueurs microsatellites possédant une grande diversité allélique devrait être suffisante pour le génotypage des poils d'ours noirs, combinée à l'utilisation de la méthode de réanalyse sélective des génotypes différents sur trois marqueurs ou moins.

Sanderlin et coll. (2012) ont récemment montré que pour une population d'ours noirs de la Géorgie, de 7 à 9 marqueurs étaient nécessaires pour obtenir une probabilité d'erreur de génotypage suffisamment faible pour la reconnaissance individuelle. Leur conclusion se base sur l'utilisation d'un algorithme récemment développé permettant de déterminer le groupe optimal de microsatellites (nombre et identité) en fonction de contraintes génétiques (probabilité d'identité [PI_{sib}], probabilité de délétion d'allèle [DA] et probabilité de faux allèles [FA] désirées) et financières (coûts des analyses génétiques $C^{(a)}$). L'algorithme de

Sanderlin et coll. (2012) se base sur une liste (L) de tous les microsatellites candidats (x_i , $i = 1$ à L) à la sélection pour lesquels différents patrons d'agencement seront testés. L'algorithme vise donc à minimiser les coûts des analyses génétiques.

$$C^{(a)} = C_0 + C_1 \sum_{i=1}^L x_i , \quad (\text{Équation 11})$$

où C_0 représente le coût de base des analyses et C_1 le coût d'analyse par échantillon. La minimisation des coûts est sujette aux contraintes suivantes :

$$PI_{sib} = \prod_{i=1}^L (PI_{sib,i} \times x_i + 1 - x_i) \leq f , \quad (\text{Équation 12})$$

où f est la limite supérieure de la probabilité PI_{sib} désirée;

$$DAM = \frac{\sum_{i=1}^L (DA_{médiane,i} \times x_i)}{\sum_{i=1}^L x_i} \leq g , \quad (\text{Équation 13})$$

où DAM est la probabilité moyenne pour tous les marqueurs de l'assemblage évalué de délétion d'allèle (DA) et où g est la limite supérieure de la probabilité DAM désirée;

$$FAM = \frac{\sum_{i=1}^L (FA_{médiane,i} \times x_i)}{\sum_{i=1}^L x_i} \leq h , \quad (\text{Équation 14})$$

où FAM est la probabilité moyenne de faux allèles (FA) pour tous les marqueurs de l'assemblage évalué et où h est la limite supérieure de la probabilité FAM désirée.

Les auteurs ont procédé à la sélection des assemblages visuellement en représentant graphiquement en 3D les résultats PI_{sib} , DAM et FAM en fonction du nombre de loci utilisés et du respect des limites supérieures f , g et h . Ils ont ainsi pu déterminer le nombre et l'identité des loci optimaux pour différencier les individus de la population à l'étude. Cette méthode permet à la fois de déterminer si le budget alloué aux analyses génétiques est réaliste par rapport au niveau de précision désiré ou de maximiser la précision en fonction du budget disponible.

Malgré le coût élevé des analyses génétiques, il est souhaitable de faire analyser le plus grand nombre possible d'échantillons de poils lors d'un inventaire. En effet, l'augmentation du nombre de recaptures fait monter rapidement le niveau de précision des estimations d'abondance et de densité. Si possible, tous les échantillons récoltés devraient donc être analysés pour maximiser les chances d'obtenir des événements de recapture. Courtois et coll. (2004) ont aussi testé plusieurs scénarios de sous-échantillonnage afin de limiter le nombre d'échantillons de poils à analyser et, ainsi, réduire le coût d'un inventaire. Tous les scénarios testés ont mené à une sous-estimation de la taille de la population (tableau 19). Si toutefois un sous-échantillonnage des poils doit être réalisé, nous proposons d'utiliser les critères suivants pour effectuer la sélection :

- Qualité des échantillons. Il est hautement préférable de conserver les échantillons avec beaucoup de poils ayant une racine. Les chances de succès sont ainsi maximisées et les risques d'erreurs, minimisés;
- Répartition des échantillons sur le fil barbelé. Puisqu'il est fort probable que deux échantillons récoltés sur des pointes de barbelés successives aient de bonnes chances d'appartenir au même ours, il devrait être suffisant d'en faire analyser un seul;
- Date et numéro de la station. Un échantillon a, en théorie, une « valeur ajoutée » plus grande s'il provient d'une station où il y avait un seul échantillon ou peu d'échantillons, comparativement à une station où il y avait de nombreux échantillons lors d'une session d'échantillonnage donnée.

Tableau 19. Estimations d'abondance obtenues en fonction de différents scénarios de sous-échantillonnage. Tiré de Courtois et coll. (2004).

	Simulation (modèle retenu dans la simulation)	Population $\pm$ S.E.
1	Avec toutes les données (M_{th} Chao)	1103 $\pm$ 108
2	Sessions d'échantillonnage 1, 2 et 3 (M_o)	627 $\pm$ 79
3	Sessions d'échantillonnage 3, 4 et 5 (M_o)	690 $\pm$ 23
4	Avec le piège 1 de chaque parcelle (M_{th})	504 $\pm$ 68
5	Avec le piège 2 de chaque parcelle (M_t Chao)	541 $\pm$ 80
6	Sessions 1 et 2 considérées comme une période de marquage alors que les sessions 3, 4 et 5 servent aux recaptures (estimateur de Petersen)	696 $\pm$ 86
7	Même scénario qu'au point précédent, mais en utilisant l'estimateur de Chapman	656 $\pm$ n/d
8	Une session de marquage (S1 et S2) avec une session de recaptures (S3 seulement) [estimateur de Petersen]	577 $\pm$ 93

6.6 Analyses démographiques

Fermeture démographique et géographique des populations d'ours noirs

La fermeture des populations est une supposition nécessaire à plusieurs analyses démographiques actuellement disponibles pour l'étude des populations. Toutefois, la plupart des études présument qu'elle est respectée sans pour autant la vérifier. Comme mentionné précédemment, le fait de considérer la population d'étude comme fermée permet d'utiliser les modèles de populations fermées, qui sont statistiquement plus robustes, pour estimer l'abondance. La condition de fermeture d'une population peut être maximisée en effectuant les inventaires à l'extérieur des périodes de mise bas ou de forte mortalité, et en délimitant l'aire d'étude en fonction de barrières naturelles. Bien que l'approche ML SECR

exige aussi que la population d'étude soit fermée, les conséquences d'une violation de cette supposition sont moins importantes étant donné la considération d'une composante spatiale dans la probabilité de détection. Nous suggérons donc l'utilisation combinée de l'approche ML SECR dans l'estimation de la densité à des méthodes de minimisation de la fermeture de la population sur le terrain.

Estimation de la densité de population chez l'ours noir

Les analyses démographiques réalisées dans le présent rapport ont permis de comparer deux approches : celle dite traditionnelle qui est toujours largement utilisée dans les études de CMR traditionnelles utilisant la reconnaissance individuelle par génotypage ($N/\hat{A}$) et une approche plus récemment développée (ML SECR), qui est censée corriger plusieurs problèmes rencontrés avec la méthode traditionnelle. L'approche ML SECR a fourni dans tous les cas des estimations de densité plus précises. Nous croyons donc que cette approche devrait être utilisée pour évaluer les populations d'ours noirs au Québec. Les ML SECR offrent divers avantages techniques et théoriques pour l'estimation de la densité des populations. En effet, l'estimation de la densité par la ML SECR peut être réalisée dans les logiciels DENSITY 4.4 et R. Ces plateformes de travail sont relativement simples et offrent une sortie de résultats facile à interpréter. La modélisation peut toutefois s'avérer longue lorsque le nombre de stations d'échantillonnage est élevé. L'extension *secr* de l'environnement de travail R semble mieux performer que DENSITY lorsque le nombre de stations est élevé. L'approche ML SECR permet de prendre en considération l'arrangement spatial des stations d'échantillonnage dans la détermination de la fonction de détection des individus. L'aire utilisée pour calculer la densité des individus est modélisée à chaque piège et elle est donc moins dépendante du positionnement des autres pièges, comme l'est l'aire effective calculée par DENSITY avec les polygones convexes et concaves. Le déplacement d'une station peut ajouter ou, au contraire, retirer une surface de l'aire effective puisque le logiciel tentera de rejoindre toutes les stations pour fermer le polygone, qu'il soit convexe ou concave. Dans la procédure de ML SECR, le déplacement d'une station a des effets beaucoup moins importants, voire nuls, sur la densité. Aussi, il est important de modéliser la probabilité de capture de tous les individus, même ceux dont la probabilité de capture est

très faible, pour éviter des biais dans l'estimation de la densité (Borchers et Efford, 2008). La probabilité de capture étant dépendante de la localisation des stations, il est important d'intégrer cette dépendance spatiale à la modélisation et à l'estimation de la densité.

Dans le présent rapport, l'intégration de la variable b aux modèles dans l'approche ML SECR a été difficile dans la majorité des cas. Il est possible que le nombre de recaptures ait été trop faible pour estimer cette variable. Nous suggérons tout de même de vérifier la performance des modèles contenant la variable b et de s'assurer que la quantité d'appât fournie est suffisamment faible pour limiter les chances d'une réponse comportementale positive.

Nous soulignons l'importance d'obtenir une estimation relativement précise de la densité de population d'ours noir. L'ours noir, comme tous les autres ursidés, est une espèce longévive dont la productivité est relativement faible. Ainsi, la période de rétablissement d'une population après une surexploitation est longue. Par exemple, Miller (1990) a montré par des simulations qu'une population d'ours noirs avec des taux de reproduction et de mortalité optimistes, et dont l'effectif a été réduit de moitié par la chasse, prendra 6 ans à se rétablir si la récolte est interdite, 9 ans si la récolte est réduite de 50 % et 17 ans si la récolte est réduite de 25 %.

7. CONCLUSION

En somme, la méthode d'inventaire utilisant la reconnaissance individuelle par génotypage des poils est réalisable et avantageuse pour l'étude des populations d'ours noirs au Québec. Nous avons déterminé un certain nombre de règles à respecter pour favoriser les chances d'éviter les principaux problèmes potentiels associés à cette approche. En résumé, nous recommandons ce qui suit :

1. Réaliser quatre sessions d'échantillonnage d'une semaine chacune à partir du mois de juin ou juillet, à l'extérieur des périodes de prélèvements par la chasse et le piégeage;

2. Utiliser un minimum de 200 stations d'échantillonnage pour obtenir une précision acceptable pour la gestion. Les stations devraient être disposées de telle sorte que chaque ours dans l'aire d'étude puisse avoir accès à au moins une station durant la période de l'inventaire. À cet effet, considérant la littérature disponible sur les domaines vitaux, nous recommandons d'utiliser une **densité minimale** d'une station tous les 25 km² dans les milieux de faible productivité (p. ex., forêt de conifères) et d'une station tous les 15 km² dans les milieux plus productifs (p. ex., forêt mixte);
3. Utiliser des enclos de fils barbelés d'une circonférence minimale de 25 m pour limiter les risques de contamination des échantillons de poils (ours multiples);
4. Récolter les poils sur les pièges avec des gants, à l'aide d'un support de couleur pâle (feuille de papier, carton, plastique) en supposant que les poils sur deux pointes de barbelés adjacentes proviennent du même ours et que ceux espacés d'une pointe vide proviennent de deux ours différents;
5. Noter, sur le terrain, le nombre de poils et la présence ou non de follicules pour chaque échantillon afin de faciliter les analyses génétiques et le sous-échantillonnage, au besoin;
6. Brûler à la torche les fils barbelés après la récolte des échantillons de poils;
7. Conserver les poils dans des enveloppes de papier pour les faire sécher à l'air libre. Lorsque l'échantillon est sec, celui-ci peut être entreposé avec des billes de silicate dans un sac de plastique (style « Ziplock ») pour le protéger de l'humidité environnante et conservé à la température de la pièce;
8. Réaliser les analyses génétiques dans les six mois suivant la récolte des poils;
9. Utiliser l'algorithme de Sanderlin et coll. (2012) lors d'une analyse génétique préliminaire afin de déterminer l'assemblage de marqueur microsatellites permettant l'identification individuelle à moindre coût;
10. Utiliser la méthode de réanalyse sélective des échantillons semblables dont les génotypes diffèrent pour trois marqueurs ou moins (3MM-pairs);
11. Utiliser l'approche ML SECR dans R pour estimer la densité de population.

8. BIBLIOGRAPHIE

- Akaike, H. 1974. « A new look at the statistical model identification ». *IEEE Transactions on Automatic Control*, vol. 19, n° 6, p. 716-723.
- Alt, G. L., G. J. Matula, Jr., F. W. Alt et J. S. Lindzey. 1980. « Dynamics of home range and movements of adult black bears in northeastern Pennsylvania ». *Bears — Their Biology and Management*, p. 131-136.
- Amstrup, S. C. et J. Beecham. 1976. « Activity patterns of radio-collared black bears in Idaho ». *The Journal of Wildlife Management*, vol. 40, n° 2, p. 340-348.
- Bacon, E. S. et M. B. Gordon. 1983. « Food preference testing of captive black bears ». *Bears — Their Biology and Management*, vol 5, p. 102-105.
- Balestrieri, A., L. Remonti, A. C. Franz, E. Capelli, M. Zenato, E. E. Dettori, F. Guidali et C. Prigioni, C. 2010. « Efficacy of passive hair-traps for genetic sampling of a low-density badger population ». *Hystrix, The Italian Journal of Mammalogy*, vol. 21, n° 2, p. 137-146.
- Bastille-Rousseau, G., D. Fortin, C. Dussault, R. Courtois, et J.-P. Ouellet. 2010. « Foraging strategies by omnivores: Are black bears actively searching for ungulate neonates or are they simply opportunistic predators? » *Ecography*, vol. 34, n° 4, p. 588-596.
- Beckmann, J. P. et J. Berger. 2003. « Rapid ecological and behavioural changes in carnivores: the responses of black bears (*Ursus americanus*) to altered food ». *Journal of Zoology*, vol. 261, n° 2, p. 207-213.
- Beecham, J. J., D. G. Reynolds et M. G. Hornocker. 1983. « Black bear denning activities and den characteristics in west-central Idaho ». *Bears — Their Biology and Management*, vol. 5, p. 79-86.
- Bellemain, E., J. E. Swenson, D. Tallmon, S. Brunberg et P. Taberlet. 2005. « Estimating population size of elusive animals with DNA from hunter-collected feces: four methods for brown bears ». *Conservation Biology*, vol. 19, n° 1, p. 150-161.
- Boersen, M. R., J. D. Clark et T. L. King. 2003. « Estimating black bear population density and genetic diversity at Tensas River, Louisiana using microsatellite DNA markers ». *Wildlife Society Bulletin*, vol. 31, n° 1, p. 197-207.

- Boileau, F. 1993. *Utilisation de l'habitat par l'ours noir (Ursus americanus) dans le parc de conservation de la Gaspésie*. Mémoire, Université Laval.
- Boileau, F., M. Crête et J. Huot. 1994. « Food habits of the black Bear, *Ursus americanus*, and habitat use in Gaspésie Park, eastern Québec ». *Canadian Field-Naturalist*, vol. 108, n° 2, p. 162-169.
- Bonin, A., E. Bellemain, P. B. Eidesen, F. Pompanon, C. Brochmann et P. Taberlet. 2004. « How to track and assess genotyping errors in population genetics studies ». *Molecular Ecology*, vol. 13, n° 11, p. 3261-3273.
- Borchers, D. L., S. T. Buckland et W. Zucchini. 2002. *Estimating Animal Abundance: Closed Populations*. Springer, Londres, 314 p.
- Borchers, D. L., et M. G. Efford. 2008. « Spatially explicit maximum likelihood methods for capture–recapture studies ». *Biometrics*, vol. 64, n° 2, p. 377-385.
- Boulanger, J., B. N. McLellan, J. G. Woods, M. F. Proctor et C. Strobeck. 2004. « Sampling design and bias in DNA-based capture-mark-recapture population and density estimates of grizzly bear ». *The Journal of Wildlife Management*, vol. 68, n° 3, p. 457-469.
- Boulanger, J., G. C. White, B. N. McLellan, J. Woods, M. Proctor et S. Himmer. 2002. « A meta-analysis of grizzly bear DNA mark-recapture projects in British Columbia, Canada: invited paper ». *Ursus*, vol. 13, p. 137-152.
- Breton, S., et F. Dufresne. 2003. *Validation des marqueurs microsatellites pour l'élaboration d'un protocole de marquage génétique chez la population d'ours noir (Ursus americanus) de la réserve faunique des Laurentides*. Université du Québec à Rimouski, pour La Société de la faune et des parcs du Québec, Direction du développement de la faune, 28 p.
- Brown, S. K., J. M. Hull, D. R. Updike, S. R. Fain et H. B. Ernest. 2009. « Black bear population genetics in California: signatures of population structure, competitive release, and historical translocation ». *Journal of Mammalogy*, vol. 90, n° 5, p. 1066-1074.
- Bunnell, F. L., et D.E.N. Tait. 1985. « Mortality Rates of North American Bears ». *Artic*, vol. 38, n° 4, p. 316-323.

- Cattet, M. R., K. Christison, N. A. Caulkett et G. B. Stenhouse. 2003. « Physiologic responses of grizzly bears to different methods of capture ». *Journal of Wildlife Diseases*, vol. 39, n° 3, p. 649-654.
- Costello, C. M. 2010. « Estimates of dispersal and home-range fidelity in American black bears ». *Journal of Mammalogy*, vol. 91, n° 1, p. 116-121.
- Coté, S. 2005. « Extirpation of a Large Black Bear Population by Introduced White-Tailed Deer ». *Conservation Biology*, vol. 19, n° 5, p. 1668–1671.
- Courtois, R., J.-P. Hamel, G. Lamontagne, R. Lemieux, J. Mercier et A. Desrosiers. 2004. *Inventaire de l'ours noir en Abitibi-Témiscamingue à l'été 2001 (premier rapport d'étape)*. Québec : ministère des Ressources naturelles, de la Faune et des Parcs, 58 p.
- Creel, S., G. Spong, J. L. Sands, J. Rotella, J. Zeigle, L. Joe, K. M. Murphy et D. Smith. 2003. « Population size estimation in Yellowstone wolves with error-prone noninvasive microsatellite genotypes ». *Molecular Ecology*, vol. 12, n° 7, p. 2003-2009.
- De Barba, M., L. P. Waits, P. Genovesi, E. Randi, R. Chirichella et E. Cetto. 2010. « Comparing opportunistic and systematic sampling methods for non-invasive genetic monitoring of a small translocated brown bear population ». *Journal of Applied Ecology*, vol. 47, n° 1, p. 172-181.
- Dice, L. R. 1938. « Some census methods for mammals ». *The Journal of Wildlife Management*, vol. 2, n° 3, p. 119-130.
- Dreher, B. P., S. R. Winterstein, K. T. Scribner, P. M. Lukacs, D. R. Etter, G. j. M. Rosa, V. A. Lopez, S. Libants et K. B. Filcek. 2007. « Noninvasive estimation of black bear population abundance incorporating genotyping errors and harvest bear ». *The Journal of Wildlife Management*, vol. 71, n° 8, p. 2684-2693.
- Drewry, J. M. 2010. *Population Abundance and Genetic Structure of Black Bears in Coastal South Carolina*. Maîtrise, University of Tennessee, 113 p.
- Efford, M. 2004. « Density estimation in live-trapping studies ». *Oikos*, vol. 106, p. 598-610.

- Efford, M. G., D. K. Dawson et C. S. Robbins. 2004. « DENSITY: software for analysing capture-recapture data from passive detector arrays ». *Animal Biodiversity and Conservation*, vol. 27.1, p. 217-228.
- Efford, M. G. 2009. DENSITY 4.4: software for spatially explicit capture-recapture. Department of Zoology, University of Otago, Dunedin, New Zealand.
- Efford, M. G. 2011. *SECR - Spatially explicit capture-recapture in R*. 25 p.
- Efford, M. G., D. L. Borchers et A. E. Byrom. 2009. « Density estimation by spatially explicit capture-recapture: likelihood-based methods ». In *Modeling demographic processes in marked populations*, Springer US (D. L. Thomson, E. G. Cooch et M. J. Conroy, éditeurs), p. 255-269.
- Foran, D. R., S. C. Minta et K. S. Heinemeyer. 1997. « DNA-based analysis of hair to identify species and individuals for population research and monitoring ». *Wildlife Society Bulletin*, vol. 25, n° 4, p. 840-847.
- Foster, R. J., et B. J. Harmsen. 2012. « A critique of density estimation from camera-trap data ». *The Journal of Wildlife Management*, vol. 76, n° 2, p. 224-236.
- Gagneux, P., C. Boesch et D. S. Woodruff. 1997. « Microsatellite scoring errors associated with noninvasive genotyping based on nuclear DNA amplified from shed hair ». *Molecular Ecology*, vol. 6, n° 9, p. 861-868.
- Gardner, B., J. A. Royle, M. T. Wegan, R. E. Rainbolt et P. D. Curtis. 2010. « Estimating black bear density using DNA data from hair snares ». *The Journal of Wildlife Management*, vol. 74, n° 2, p. 318-325.
- Garshelis, D. L. 2011. « Andean bear density and abundance estimates: how reliable and useful are they? » *Ursus*, vol. 22, n° 1, p. 47-64.
- Garshelis, D. L., et K. V. Noyce. 2008. *Ecology and Population Dynamics of Black Bears in Minnesota*. Minnesota Department of Natural Resources, p. 356-362.
- Garshelis, D. L., H. B. Quigley, C. R. Villarrubia et M. R. Pelton. 1983. « Diel movements of black bears in the southern Appalachians ». *Bears — Their Biology and Management*, vol. 5, p. 11-19.
- Gignac, L., V. Brodeur, C. Daigle, S. Lefort, F. Landry et J. Bouchard. 2008. *Gros gibier au Québec. Données de récolte 1^{er} mai 2006 au 30 avril 2007*. Québec : ministère des Ressources naturelles et de la Faune, 55 p.

- Goossens, B., L. P. Waits et P. Taberlet. 1998. « Plucked hair samples as a source of DNA: reliability of dinucleotide microsatellite genotyping ». *Molecular Ecology*, vol. 7, n° 9, p. 1237-1241.
- Graber, D. M., et M. White. 1983. « Black bear food habits in Yosemite National Park ». *Bears — Their Biology and Management*, vol. 5, p. 1-10.
- Greenwood, P. J. 1980. « Mating systems, philopatry and dispersal in birds and mammals ». *Animal Behaviour*, vol. 28, n° 4, p. 1140-1162.
- Grogan, R. G., et F. G. Lindzey. 1999. « Estimating population size of a low-density black bear population using capture-resight ». *Ursus*, vol. 11, p. 117-122.
- Harris, R. B., J. Winnie, S. J. Amish, A. Beja-Pereira, R. Godinho, V. Costa et G. Luikart. 2010. « Argali abundance in the afghan Pamir using capture–recapture modeling from fecal DNA ». *The Journal of Wildlife Management*, vol. 74, n° 4, p. 668-677.
- Hellgren, E. C., et M. R. Vaughan. 1989. « Demographic analysis of a black bear population in the great dismal swamp ». *The Journal of Wildlife Management*, vol. 53, n° 4, p. 969-977.
- Herrero, S. 1983. « Social behaviour of black bears at a garbage dump in Jasper National Park. *Bears — Their Biology and Management*, vol. 5, p. 54-70.
- Higuchi, R., C. H. von Beroldingen, G. F. Sensabaugh et H. A. Erlich. 1988. « DNA typing from single hairs ». *Nature*, vol. 332, p. 543-546.
- Immell, D., et R. G. Anthony. 2008. « Estimation of black bear abundance using a discrete DNA sampling device ». *The Journal of Wildlife Management*, vol. 72, n° 1, p. 324-330.
- Jarne, P., et P. J. L. Lagoda. 1996. « Microsatellites, from molecules to populations and back ». *Trends in Ecology & Evolution*, vol. 11, n° 10, p. 424-429.
- Jett, D. A., et J. D. Nichols. 1987. « A field comparison of nested grid and trapping web density estimators ». *Journal of Mammalogy*, vol. 68, n° 4, p. 888-892.
- Johnson, J. B., et K. S. Omland. 2004. « Model selection in ecology and evolution ». *Trends in Ecology & Evolution*, vol. 19, n° 2, p. 101-108.
- Jolicoeur, H. 1996. *Le marquage à la tétracycline comme méthode pour estimer les populations d'ours noir sur de grandes superficies*. Québec : ministère de l'Environnement et de la Faune, 22 p.

- Jolicoeur, H. 1997a. *Bilan de l'exploitation de l'ours noir au Québec. III-Diagnostic sur l'exploitation de l'ours noir. Période 1984-1995*. Québec : ministère de l'Environnement et de la Faune, 50 p.
- Jolicoeur, H. 1997b. *Démarche pour analyser et interpréter les statistiques de récolte d'ours noir*. Québec : ministère de l'Environnement et de la Faune, 72 p.
- Jolicoeur, H. 2004. *Estimation de la densité d'ours noirs dans différents types de végétation à l'aide de traceurs radioactifs. Période 1984-1994*. Québec : ministère des Ressources naturelles, de la Faune et des Parcs, 52 p.
- Jolicoeur, H., et G. Lamontagne. 1997. *Bilan de l'exploitation de l'ours noir au Québec. I-Portrait de la récolte d'ours noirs. Période 1984-1995*. Québec : ministère de l'Environnement et de la Faune, 53 p.
- Jolicoeur, H., et R. Lemieux. 1984. *Évaluation de la densité de l'ours noir de la réserve de Papineau-Labelle au moyen d'un traceur radioactif*. Québec : ministère du Loisir, de la Chasse et de la Pêche, 34 p.
- Jolicoeur, H., et R. Lemieux. 1990. *Comparaison de deux méthodes pour évaluer la densité de l'ours noir à la réserve Papineau-Labelle*. Québec : ministère du Loisir, de la Chasse et de la Pêche, 78 p.
- Jolicoeur, H., et R. Lemieux. 1994. *Quelques aspects de la reproduction de l'ours noir au Québec*. Québec : ministère de l'Environnement et de la Faune, 67 p.
- Kalinowski, S. T., M. L. Taper et T. C. Marshall. 2007. « Revising how the computer program CERVUS accommodates genotyping error increases success in paternity assignment ». *Molecular Ecology*, vol. 16, n° 5, p. 1099-1006.
- Kasworm, W. F., M. F. Proctor, C. Servheen et D. Paetkau. 2007. « Success of grizzly bear population augmentation in northwest Montana ». *The Journal of Wildlife Management*, vol. 71, n° 4, p. 1261-1266.
- Kemp, G. A. 1976. « The dynamics and regulation of black bear, *Ursus americanus*, population in northern Alberta. Third International Conference on Bear ». *Bear — Their Biology and Management*, vol. 40, p. 191-197.
- Knapp, S. M., B. A. Craig et L. P. Waits. 2009. « Incorporating genotyping error into non-invasive DNA-based mark-recapture population estimates ». *The Journal of Wildlife Management*, vol. 73, n° 4, p. 598-604.

- Lamontagne, G., H. Jolicoeur et S. Lefort. 2006. *Plan de gestion de l'ours noir 2006-2013*. Québec : ministère des Ressources naturelles et de la Faune, 487 p.
- Larivière, S. 2001. « *Ursus americanus* ». *Mammalian Species*, vol. 647, p. 1-11.
- Leblanc, N., et J. Huot. 2000. *Sélection d'habitats et utilisation du milieu par l'ours noir (Ursus americanus) dans une aire protégée de dimension restreinte : le parc national de Forillon*. Maîtrise, Université Laval, 115 p.
- Leblanc, N., et C. Samson. 2006. *Paramètre d'exposition chez les mammifères – Ours noir. Fiche descriptive*. Centre d'expertise en analyse environnementale du Québec, ministère du Développement durable, de l'Environnement et des Parcs, 17 p.
- Lebreton, J.-D., K. P. Burnham, J. Clobert et D. R. Anderson. 1992. « Modeling survival and testing biological hypotheses using marked animals: a unified approach with case studies ». *Ecological Monographs*, vol. 62, n° 1, p. 67-118.
- Lecount, A. L. 1982. « Characteristics of a central arizona black bear population ». *The Journal of Wildlife Management*, vol. 46, n° 4, p. 861-868.
- Lemieux, R. 1985. *Résultat des opérations de marquage de l'ours noir dans la réserve de Papineau-Labelle dans le cadre de l'étude sur l'évaluation de la densité de l'ours noir au moyen d'un traceur radioactif*. Québec : ministère du Loisir, de la Chasse et de la Pêche, 33 p.
- Lemieux, R., et A. Desrosiers. 2002. *Description des techniques développées pour la collecte des poils d'ours noir, Ursus americanus, et l'ingestion de la tétracycline dans la réserve faunique des Laurentides à l'été 2000*. Société de la Faune et des Parcs du Québec, 38 p.
- Lindahl, T. 1993. « Instability and decay of the primary structure of DNA ». *Nature*, vol. 362, p. 709-715.
- Lukacs, P. M., et K. P. Burnham. 2005. « Review of capture–recapture methods applicable to noninvasive genetic sampling ». *Molecular Ecology*, vol. 14, n° 13, p. 3909-3919.
- Mainguy, J., et L. Bernatchez. 2007. « L'analyse de l'ADN sans manipulation des animaux: un outil incontournable pour la gestion et la conservation des espèces rares ou élusives ». *Le Naturaliste canadien*, vol. 131, n° 1, p. 51-58.

- Massé, S. 2012. *Impacts d'un site de nourrissage à des fins d'observation sur l'utilisation de l'espace et la sélection de l'habitat de l'ours noir (Ursus americanus) en forêt boréale*. Mémoire, Université du Québec à Chicoutimi.
- Mattson, D. J. 1990. « Human impacts on bear habitat use ». Eighth International Conference on Bear Research and Management, International Association for Bear Research and Management, Victoria, Colombie-Britannique, p. 33-56.
- McCall, B. S. 2002. *Noninvasive genetic sampling reveals black bear population dynamics driven by changes in food productivity*. Maîtrise, University of Idaho, 51 p.
- McDonald, T. L., S. C. Amstrup et B. F. J. Manly. 2003. « Tag loss can bias Jolly-Seber capture-recapture estimates ». *Wildlife Society Bulletin*, vol. 31, n° 3, p. 814-822.
- McKelvey, K. S., et M. K. Schwartz. 2004. « Genetic errors associated with population estimation using non-invasive molecular tagging: problems and new solutions ». *The Journal of Wildlife Management*, vol. 68, n° 3, p. 439-448.
- McKelvey, K. S. et M. K. Schwartz. 2005. « Dropout: a program to identify problem loci and samples for noninvasive genetic samples in a capture-mark-recapture framework ». *Molecular Ecology Notes*, vol. 5, n° 3, p. 716-718.
- McLaughlin, C. R., G. J. Matula, Jr. et R. J. O'Connor. 1994. « Synchronous reproduction by Maine black bears ». *Bears – Their Biology and Management*, vol. 9, p. 471-479.
- Meredith, E., J. Rodzen, J. Banks et K. Jones. 2009. « Characterization of 29 tetranucleotide microsatellite loci in black bear (*Ursus americanus*) for use in forensic and population applications ». *Conservation Genetics*, vol. 10, n° 3, p. 693-696.
- Miller, C. R., P. Joyce et L. P. Waits. 2005. « A new method for estimating the size of small populations from genetic mark-recapture data ». *Molecular Ecology*, vol. 14, n° 7, p. 1991-2005.
- Miller, S. D. 1990. « Population management of bears in North America ». *Bears – Their Biology and Management*, vol. 8, p. 357-373.
- Mills, S. L., J. J. Citta, K. P. Lair, M. K. Schwartz D. A. Tallmon. 2000. « Estimating animal abundance using noninvasive DNA sampling: promise and pitfalls ». *Ecological Applications*, vol. 10, n° 1, p. 283-294.
- Ministère des Ressources naturelles et de la Faune du Québec. 2010. *La chasse sportive au Québec – Principale règles, Du 1^{er} avril 2010 au 31 mars 2012*.

- Mowat, G., D. C. Heard, D. R. Seip, K. G. Poole, G. Stenhouse et D. W. Paetkau. 2004. « Grizzly *Ursus arctos* and black bear *U. americanus* densities in the interior mountains of North America ». *Wildlife Biology*, vol. 11, n° 1, p. 31-48.
- Navidi, W., N. Arnheim et M. S. Waterman. 1992. « A multiple-tube approach for accurate genotyping of very small DNA samples bu using PCR: statistical considerations ». *American Journal of Human Genetics*, vol. 50, n° 2, p. 347-359.
- Noyce, K. V., P. B. Kannowski et M. R. Riggs. 1997. « Black bears as ant-eaters: seasonal associations between bear myrmecophagy and ant ecology in north-central Minnesota ». *Canadian Journal of Zoology*, vol. 75, n° 10, p. 1671-1686.
- Obbard, M. E., E. J. Howe et C. J. Kyle. 2010. « Empirical comparison of density estimators for large carnivores ». *Journal of Applied Ecology*, vol. 47, n° 1, p. 76-84.
- Otis, D. L., K. P. Burnham, G. C. White et D. R. Anderson. 1978. « Statistical inference from capture data on closed animal population ». *Wildlife Monographs*, vol. 62, p. 3-135.
- Pacas, C. J., et P. C. Paquet. 1994. « Analysis of black bear home range using a geographic information system ». *Bears – Their Biology and Management*, vol. 9, p. 419-425.
- Paetkau, D. 2003. « An empirical exploration of data quality in DNA-based population inventories ». *Molecular Ecology*, vol. 12, n° 6, p. 1375-1387.
- Paetkau, D. 2004. « The optimal number of markers in genetic cpature-mark-recapture studies ». *The Journal of Wildlife Management*, vol. 68, n° 3, p. 449-452.
- Paetkau, D., et C. Strobeck. 1994. « Microsatellite analysis of genetic variation in black bear populations ». *Molecular Ecology*, vol. 3, n° 5, p. 489-495.
- Palsboll, P. J., et J. Allen. 1997. « Genetic tagging of humpback whales ». *Nature*, vol. 388, p. 767.
- Panasci, M., W. B. Ballard, S. Breck, D. Rodriguez, L. D. Densmore III, D. B. Wester et R. J. Baker. 2011. « Evaluation of fecal DNA preservation techniques and effects of sample age and diet on genotyping success ». *The Journal of Wildlife Management*, vol. 75, n° 7, p. 1616-1624.

- Peacock, E., K. Titus, D. L. Garshelis, M. M. Peacock et M. Kuc. 2011. « Mark-recapture using tetracycline and genetics reveal record-high bear density ». *The Journal of Wildlife Management*, vol. 75, n° 6, p. 1513-1520.
- Pearse, D. E., M. Eckerman, F. J. Janzen et J. C. Avise. 2001. « A genetic analogue of "mark-recapture" methods for estimating population size: an approach based on molecular parentage assessments ». *Molecular Ecology*, vol. 10, n° 11, p. 2711-2718.
- Pledger, S. 2000. « Unified maximum likelihood estimates for closed capture-recapture models using mixtures ». *Biometrics*, vol. 56, n° 2, p. 434-442.
- Pollock, K. H., J. D. Nichols, C. Brownie et J. E. Hines. 1990. « Statistical inference for capture-recapture experiments ». *Wildlife Monographs*, vol. 107, p. 3-97.
- Poole, K. G., G. Mowat et D. A. Fear. 2001. « DNA-based population estimate for grizzly bears *Ursus arctos* in northeastern British Columbia, Canada ». *Wildlife Biology*, vol. 7, n° 2, p. 105-115.
- Proctor, M., B. N. McLellan, J. Boulanger, C. Apps, G. Stenhouse, D. Paetkau et G. Mowat. 2010. « Ecological investigations of grizzly bears in Canada using DNA from hair, 1995-2005: a review of method and progress ». *Ursus*, vol. 21, p. 169-188.
- R Development Core Team. 2011. *R: A language and environment for statistical Computing*. R Foundation for Statistical Computing. Vienne, Autriche.
- Rexstad, E., et K. P. Burnham. 1992. *User's guide for interactive program CAPTURE*. Institute of Arctic Biology and Department of Biology and Wildlife, University of Alaska.
- Robinson, S. J., L. P. Waits et I. D. Martin. 2007. « Evaluating population structure of black bears on the Kenai peninsula using mitochondrial and nuclear DNA analyses ». *Journal of Mammalogy*, vol. 88, n° 5, p. 1299-2007.
- Robinson, S. J., L. P. Waits et I. D. Martin. 2009. « Estimating abundance of American black bears using DNA-based capture-mark-recapture models ». *Ursus*, vol. 20, p. 1-11.
- Rogers, L. L. 1987. « Effects of food supply and kinship on social behavior, movements, and population growth of black bears in northeastern Minnesota ». *Wildlife Monographs*, p. 3-72.

- Romain-Bondi, K. A., R. B. Wielgus, L. Waits, W. F. Kasworm, M. Austin et W. Wakkinen. 2004. « Density and population size estimates for North Cascade grizzly bears using DNA hair-sampling techniques ». *Biological Conservation*, vol. 117, p. 417-428.
- Roon, D. A., L. P. Waits et K. C. Kendall. 2003. « A quantitative evaluation of two methods for preserving hair samples ». *Molecular Ecology Notes*, vol. 3, n° 1, p. 163-166.
- Roon, D. A., L. P. Waits et K. C. Kendall. 2005. « A simulation test of the effectiveness of several methods for error-checking non-invasive genetic data ». *Animal Conservation*, vol. 8, n° 2, p. 203-215.
- Roy, J., V. Albert et L. Bernatchez. 2007. *Projet d'inventaire de l'ours noir dans la zone 10 par la technique de capture-recapture à l'aide de marqueurs génétiques (Projet Outaouais 2005)*. Université Laval et ministère des Ressources naturelles et de la Faune du Québec, 164 p.
- Roy, J., G. Yannic, S. Côté et L. Bernatchez. 2012. « Negative density-dependent dispersal in the American black bear (*Ursus americanus*) revealed by noninvasive sampling and genotyping ». *Ecology and Evolution*, vol. 2, n° 3, p. 525-537.
- Saiki, R. K., D. H. Gelfand, S. Stoffel, S. J. Scharf, R. Higuchi, G. T. Horn, K. B. Mullis et H. A. Erlich. 1988. « Primer-Directed Enzymatic Amplification of DNA with a Thermostable DNA Polymerase ». *Science*, vol. 239, p. 487-491.
- Samson, C. 1996. *Modèle d'indice de qualité de l'habitat pour l'ours noir (Ursus americanus) au Québec*. Québec : ministère de l'Environnement et de la Faune, 57 p.
- Samson, C., et J. Huot. 1994. *Écologie et dynamique de la population d'ours noir (Ursus americanus) du parc national de la Mauricie. Rapport final*. Université Laval, 214 p.
- Sanderlin, J. S., B. C. Faircloth, B. Shamblin et M. J. Conroy. 2009. « Tetranucleotide microsatellite loci from the black bear (*Ursus americanus*) ». *Molecular Ecology Resources*, vol. 9, n° 1, p. 288-291.
- Sanderlin, J. S., N. Lazar, M. J. Conroy et J. Reeves. 2012. « Cost-efficient selection of a marker panel in genetic studies ». *The Journal of Wildlife Management*, vol. 76, n° 1, p. 88-94.

- Schwartz, C. C., et A. W. Franzmann. 1991. « Interrelationship of black bears to moose and forest succession in the northern coniferous forest ». *Wildlife Monographs*, vol. 113, p. 1-58.
- Schwarz, C. J., et A. N. Arnason. 1996. « A general methodology for the analysis of capture-recapture experiments in open populations ». *Biometrics*, vol. 52, p. 860-873.
- Seber, G. A .F. 2002. *The Estimation of Animal Abundance and Related Parameters*. Blackburn Press, 654 p.
- Settlage, K. E., F. T. Van Manen, J. D. Clark et T. L. King. 2008. « Challenges of DNA-based mark–recapture studies of American black bears ». *The Journal of Wildlife Management*, vol. 72, n° 4, p. 1035-1042.
- Sloane, M. A., P. Sunnucks, D. Alpers, L. B. Beheregaray et A. C. Taylor. 2000. « Highly reliable genetic identification of individual northern hairy-nosed wombats from single remotely collected hairs: a feasible censusing method ». *Molecular Ecology*, vol. 9, n° 9, p. 1233-1240.
- Solberg, K. H., E. Bellemain, O.-M. Drageset, P. Taberlet et J. E. Swenson. 2006. « An evaluation of field and non-invasive genetic methods to estimate brown bear (*Ursus arctos*) population size ». *Biological Conservation*, vol. 128, n° 2, p. 158-168.
- Stanley, T. R., et K. P. Burnham. 1999. « A closure test for time-specific capture-recapture data ». *Environmental and Ecological Statistics*, vol. 6, p. 197-209.
- Stanley, T. R., et D. R. Jon. 2005. « A program for testing capture-recapture data for closure ». *Wildlife Society Bulletin*, vol. 33, n° 2, p. 782-785.
- Stanley, T. R., et J. D. Richards. 2004. *CloseTest: a program for testing capture-recapture data for closure (software manual)*. U.S. Geological Survey, Fort Collins Science Center.
- Stanley, T. R., et J. D. Richards. 2005. « Software review: a program for testing capture–recapture data for closure ». *Wildlife Society Bulletin*, vol. 33, n° 2, p. 782-785.
- Taberlet, P., et J. Bouvet. 1992. « Bear conservation genetics ». *Nature*, vol. 358, p. 197.
- Taberlet, P., J.-J. Camarra, S. Griffin, E. Uhres, O. Hanotte, L. P. Waits, C. Dubois-Paganon, T. Burke et J. Bouvet. 1997. « Noninvasive genetic tracking of the

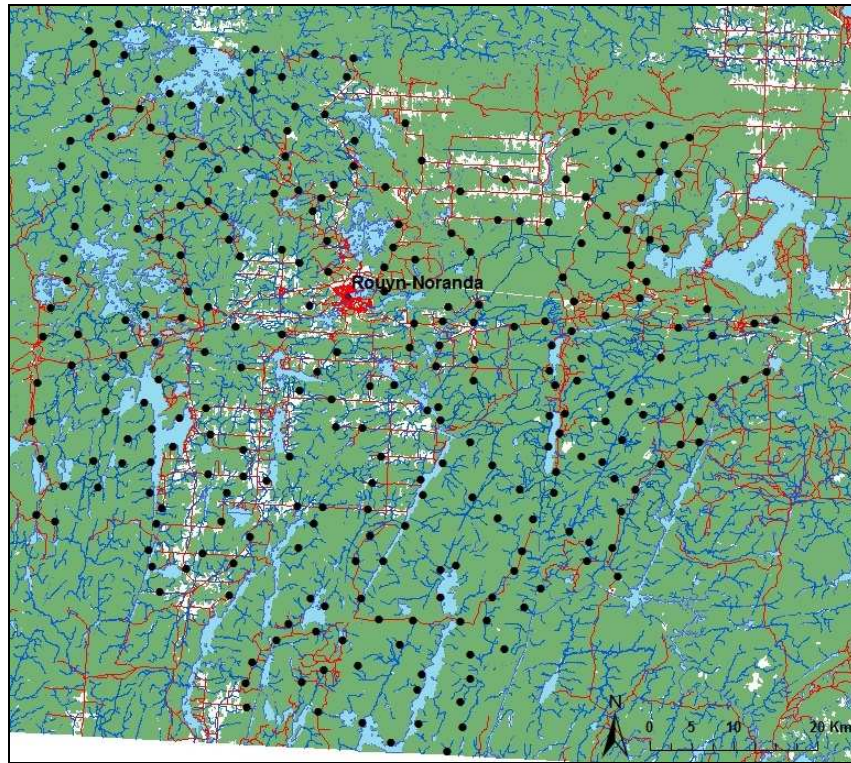
- endangered Pyrenean brown bear population ». *Molecular Ecology*, vol. 6, n° 9, p. 869-876.
- Taberlet, P., S. Griffin, B. Goossens, S. Questiau, V. Manceau, N. Escaravage, L. P. Waits et J. Bouvet. 1996. « Reliable genotyping of samples with very low DNA quantities using PCR ». *Nucleic Acids Research*, vol. 24, n° 16, p. 3189-3194.
- Taberlet, P., et G. Luikart. 1999. « Non-invasive genetic sampling and individual identification ». *Biological Journal of the Linnean Society*, vol. 68, n° 1-2, p. 41-55.
- Taberlet, P., H. Mattock, C. Dubois-Paganon et J. Bouvet. 1993. « Sexing free-ranging brown bears *Ursus arctos* using hairs found in the field ». *Molecular Ecology*, vol. 2, n° 6, p. 399-403.
- Taberlet, P., L. P. Waits et G. Luikart. 1999. « Noninvasive genetic sampling: look before you leap ». *Trends in Ecology & Evolution*, vol. 14, n° 8, p. 323-327.
- Tredick, C. A., M. R. Vaughan, D. F. Stauffer, S. L. Simek et T. Eason, T. 2007. « Sub-sampling genetic data to estimate black bear population size: a case study ». *Ursus*, vol. 18, n° 2, p. 179-188.
- Triant, D. A., R. M. Pace et M. Stine, M. 2004. « Abundance, genetic diversity and conservation of Louisiana black bears (*Ursus americanus luteolus*) as detected through noninvasive sampling ». *Conservation Genetics*, vol. 5, p. 647-659.
- Wagner, H. W., et K. M. Sefc. 1999. *IDENTITY 1.0*. Center for Applied Genetics, University of Agricultural Sciences, Vienne, Autriche.
- Waits, J. L., et P. L. Leberg. 2000. « Biases associated with population estimation using molecular tagging ». *Animal Conservation*, vol. 3, n° 3, p. 191-199.
- Waits, L., D. Paetkau et C. Strobeck. 1999. « Genetics of the bears of the world ». *Bear Conservation Act* (Éd. C. Servheen), p. 25-32.
- Waits, L. P., G. Luikart et P. Taberlet. 2001. « Estimating the probability of identity among genotypes in natural populations: cautions and guidelines ». *Molecular Ecology*, vol. 10, n° 1, p. 249-256.
- Waits, L. P., et D. Paetkau. 2005. « Noninvasive genetic sampling tools for wildlife biologists: a review of applications and recommendations for accurate data collection ». *The Journal of Wildlife Management*, vol. 69, n° 4, p. 1419-1433.

- Waser, P. M., et W. T. Jones. 1983. « Natal philopatry among solitary mammals ». *The Quarterly Review of Biology*, vol. 58, n° 3, p. 355-390.
- Weaver, K. M., et M. R. Pelton. 1994. « Denning ecology of black bears in the Tensas River basin of Louisiana ». *International Conference on Bear Resources and Management*, vol. 9, n° 1, p. 427-433.
- White, G. C., D. R. Anderson, K. P. Burnham et D. L. Otis. 1982. *Capture-recapture and removal methods for sampling closed populations*. Los Alamos National Laboratory Publication, Nouveau-Mexique, États-Unis, 235 p.
- White, G. C., et K. P. Burnham. 1999. « Program MARK: Survival estimation from populations of marked animals ». *Bird Study*, vol. 46, p. 120-128.
- Wilson, K. R., et D. R. Anderson. 1985. « Evaluation of two density estimators of small mammal population size ». *Journal of Mammalogy*, vol. 66, n° 1, p. 13-21.
- Wimsatt, W. A. 1963. « Delayed implantation in the Ursidae with particular reference to the black bear (*Ursus americanus* Pallas) ». Pages 49-76 dans A. C. Enders (éditeur), *Delayed Implantation*, Univ. Chicago Press, Chicago, IL.
- Woods, J. G., D. Paetkau, D. Lewis, B. N. McLellan, M. Proctor et C. Strobeck. 1999. « Genetic tagging of free-ranging black and brown bears ». *Wildlife Society Bulletin*, vol. 27, n° 3, p. 616-627.
- Wright, J. A., R. J. Barker, M. R. Schofield, A. C. Frantz, A. E. Byrom et D. M. Gleeson. 2009. « Incorporating genotype uncertainty into mark-recapture-type models for estimating abundance using DNA samples ». *Biometrics*, vol. 65, n° 3, p. 833-840.
- Yamamoto, K., T. Tsubota, T. Komatsu, A. Katayama, T. Murase, I. Kita et T. Kudo. 2002. « Sex identification of Japanese black bear, *Ursus thibetanus japonicus*, by PCR based on amelogenin gene ». *Journal of Veterinary Medical Science*, vol. 64, n° 6, p. 505-508.

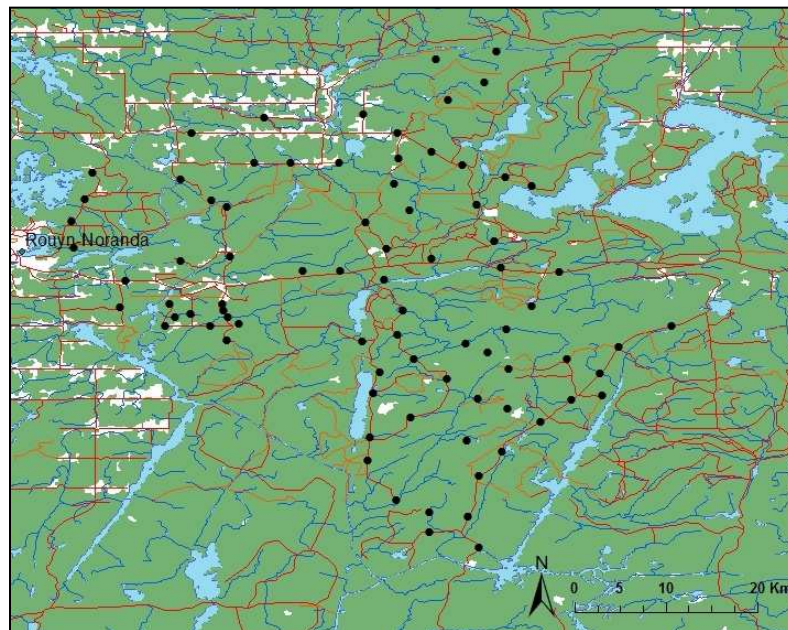
Annexe 1. Modalité d'échantillonnage et résultats partiels d'études de CMR utilisant la reconnaissance individuelle par génotypage des poils réalisées chez l'ours noir.

Aire d'étude (km ²)	Aire des parcelles (km ²)	Nombre de parcelles	Statut des parcelles	Durée et période de l'échantillonnage	Nombre d'individus identifiés	Estimation de la population	Densité (ours/10 km ²)	Domaine	
								vital d'une femelle (km ²)	Référence
36 848	> 25,0	239	Fixes	5 sessions (5-7jrs/ses.), 22 juin au 26 juillet	165	1922	n.d.	131	Deher et coll., 2007
3 000	25,0	120	Fixes	4 sessions (7 jrs/ses.), 11 juillet au 4 août	151	2077	0,96	45	Roy et coll., 2007
50	1,5	33	Fixes	8 sessions (7 jrs/ses.), de juin à août	85	128	n.d.	2	Tredick et coll., 2007
967	19,0	51	Fixes	8 sessions (7 jrs/ses.), de juin à août	25	39	n.d.	24	Tredick et coll., 2007
160	2,5	65	Fixes	10 sessions (7 jrs/ses.)	129	n.d.	n.d.	8	Settlage et coll., 2008
329	5,8	57	Fixes	10 sessions (7 jrs/ses.)	60	n.d.	n.d.	14	Settlage et coll., 2008
157	7,1	38	Fixes	8 sessions (7 jrs/ses.), juin et juillet	47	114	2,0	14	Gardner et coll., 2010
112	1,0	113	Fixes	3 sessions (21 jrs/ses.), 1 ^{er} juin au 1 ^{er} août	32	46	1,8	14	Immell et coll., 2008
137	1,0	123	Fixes	3 sessions (21 jrs/ses.), 1 ^{er} juin au 1 ^{er} août	30	57	2,0	14	Immell et coll., 2008
155	1,0	131	Fixes	3 sessions (21jrs/ses.), 1 ^{er} juin au 1 ^{er} août	47	67	2,3	14	Immell et coll., 2008
145	1,0	125	Fixes	3 sessions (21 jrs/ses.), 1 ^{er} juin au 1 ^{er} août	46	65	2,1	14	Immell et coll., 2008
329	2,7	122	Fixes	14 sessions (7 jrs/ses.), 27 juillet au 2 novembre	58	119	3,6	24	Boersen et coll., 2003

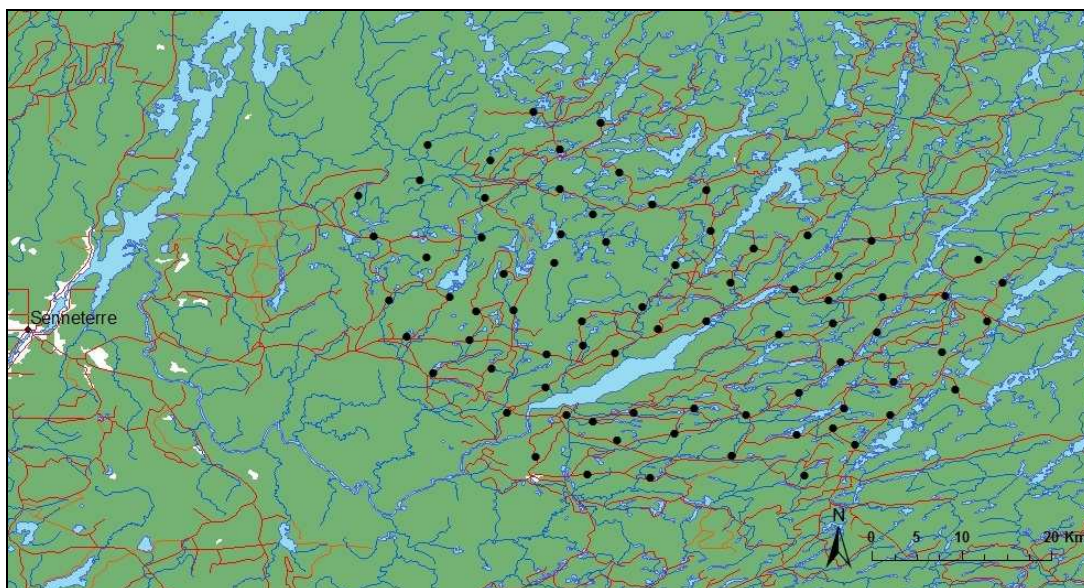
Annexe 2. Cartes des secteurs d'inventaire de l'ours noir en Abitibi-Témiscamingue
de 2001 à 2003.



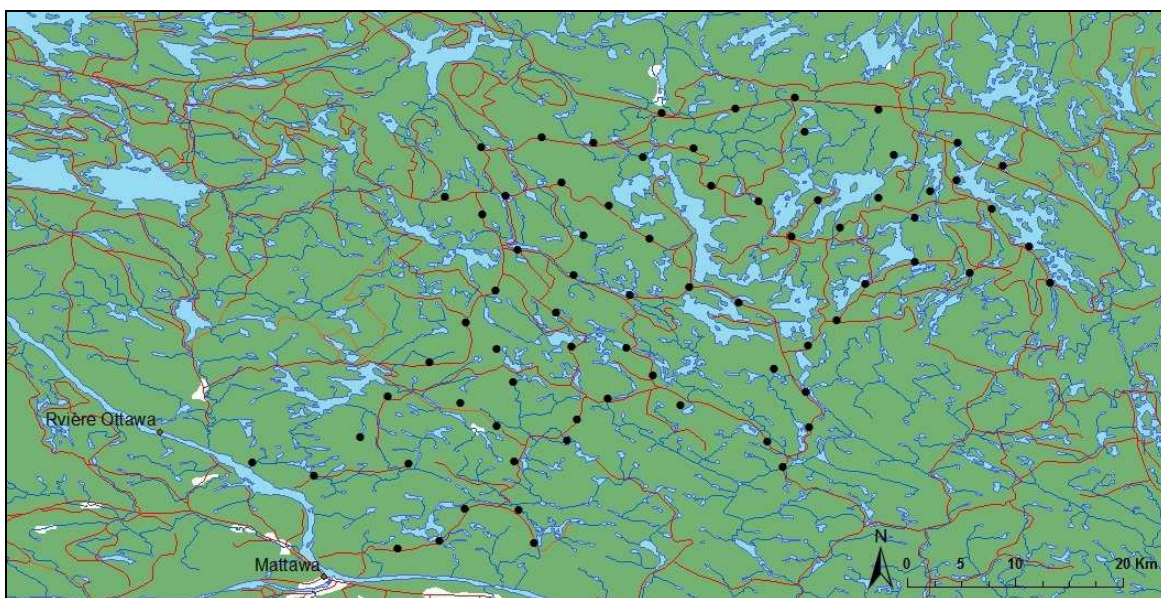
Disposition spatiale des stations d'échantillonnage lors de l'étude de l'été 2001, en périphérie de Rouyn-Noranda dans la région de l'Abitibi-Témiscamingue



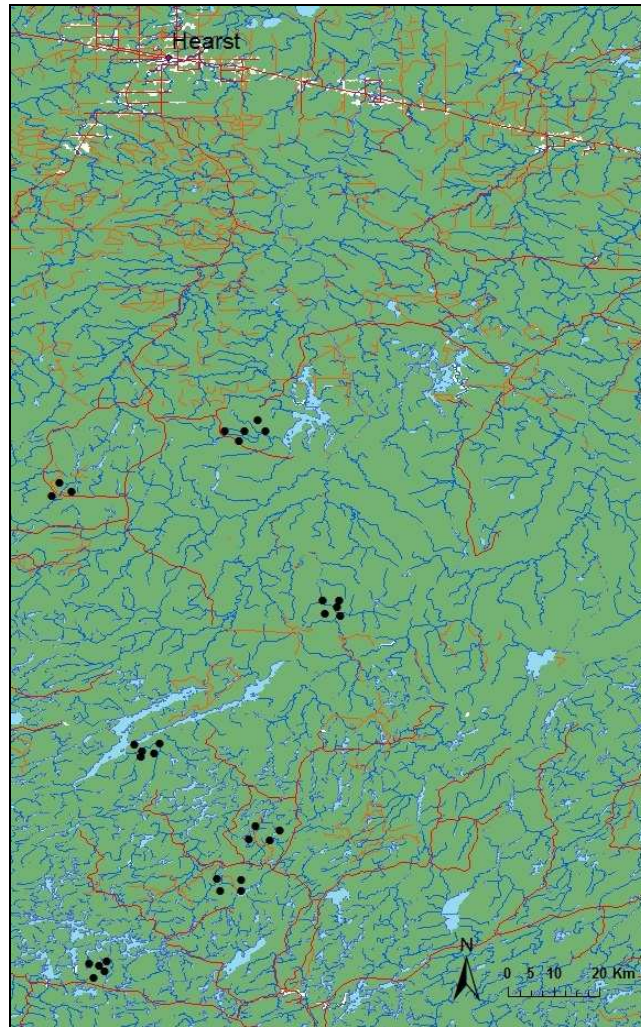
Disposition spatiale des stations d'échantillonnage lors de l'étude de l'été 2002 à Roger-Rapide 7 (R2002)



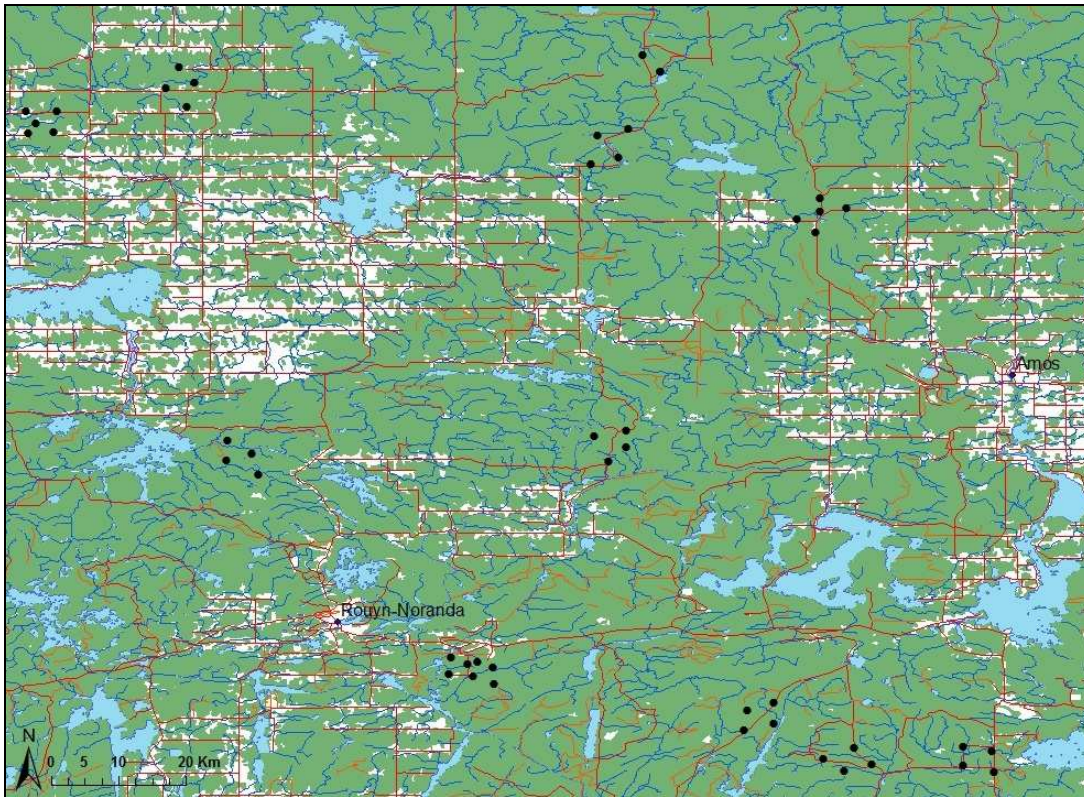
Disposition spatiale des stations d'échantillonnage lors de l'étude de l'été 2002 à Senneterre (S2002)



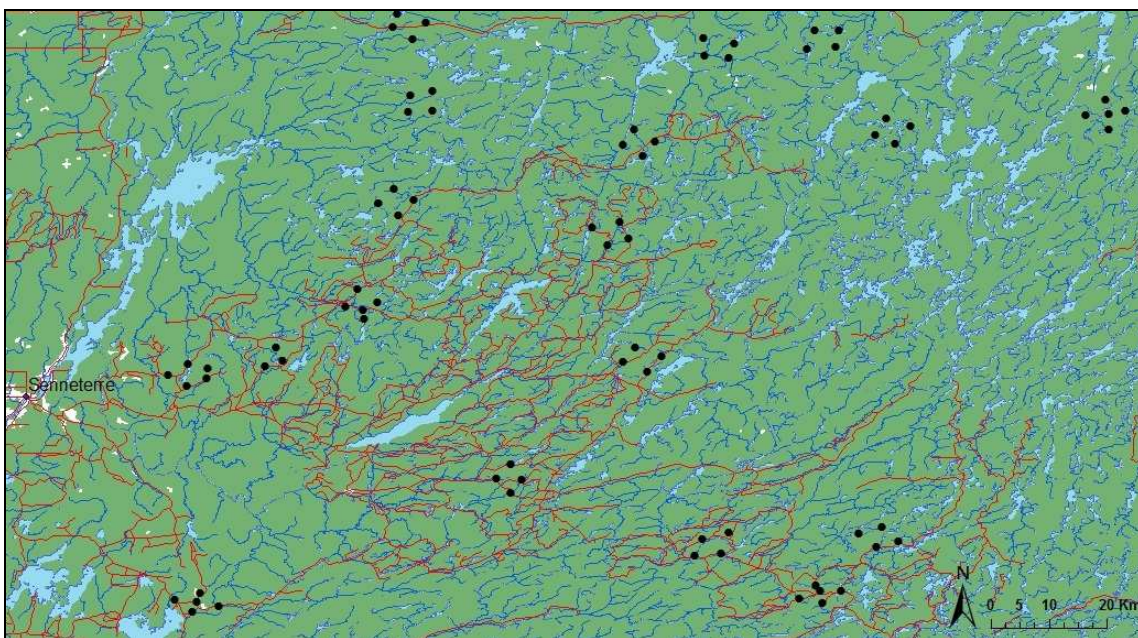
Disposition spatiale des stations d'échantillonnage lors de l'étude de l'été 2002 à Témiscamingue (T2002)



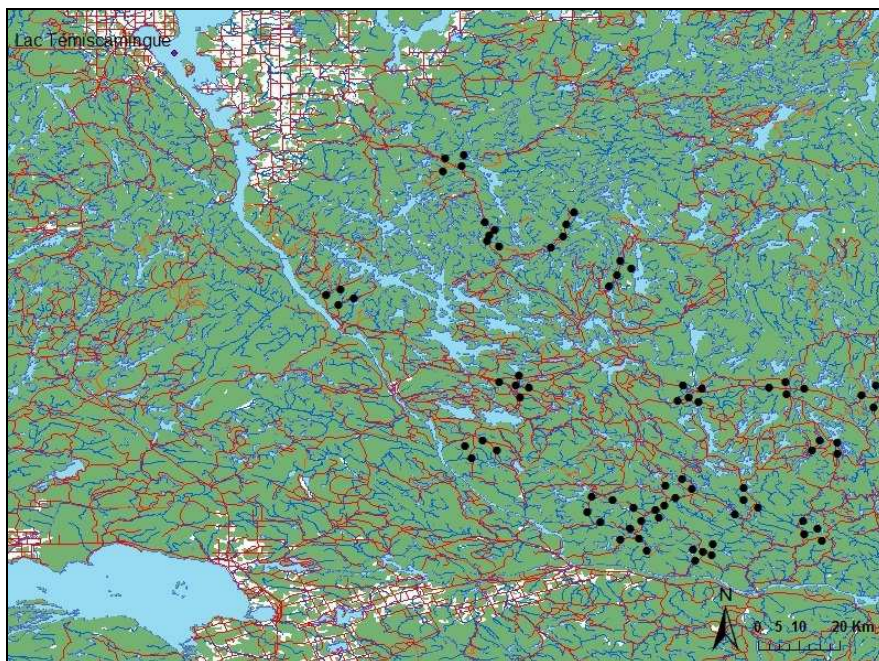
Disposition spatiale des stations d'échantillonnage lors de l'étude de l'été 2003 à Rouyn-Noranda Ouest (R2003-O)



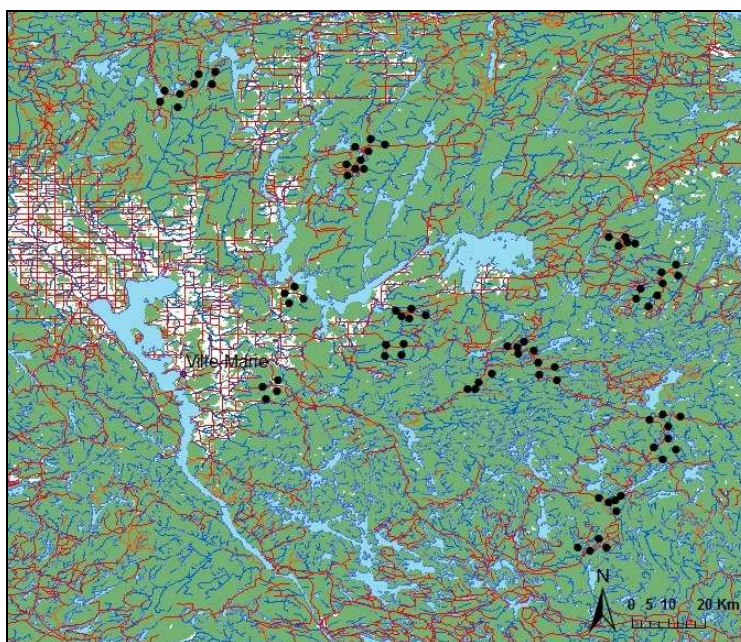
Disposition spatiale des stations d'échantillonnage lors de l'étude de l'été 2003 à Rouyn-Noranda Est (R2003-E)



Disposition spatiale des stations d'échantillonnage lors de l'étude de l'été 2003 à Senneterre (S2003)



Disposition spatiale des stations d'échantillonnage lors de l'étude de l'été 2003 à Témiscamingue (T2003)



Disposition spatiale des stations d'échantillonnage lors de l'étude de l'été 2003 à Ville-Marie (V2003)

Annexe 3. Mode d'emploi pour le logiciel IDENTITY 4.

Le logiciel IDENTITY 4 permet de déterminer les paires de génotypes semblables parmi un lot de génotypes identifiés sur au moins deux marqueurs microsatellites. Une description plus détaillée du programme et les instructions d'utilisation sont disponibles sur le site <http://www.uni-graz.at/~sefck/>.

1. Le programme IDENTITY 4 est disponible au <http://www.uni-graz.at/~sefck/>. Cliquez sur *Download*. Il est préférable d'enregistrer le logiciel directement sur le C:\. Créez un dossier que vous nommerez IDENTITY et dans lequel vous déposerez le logiciel. Il sera plus facile par la suite d'ouvrir le logiciel à l'aide d'une fenêtre DOS.
2. Ouvrez une fenêtre DOS. La route d'accès pour la fenêtre DOS est différente d'un système Windows à l'autre, mais elle est disponible, la plupart du temps, dans les accessoires du menu « Démarrer ».
3. Une fois la fenêtre DOS ouverte, il faut « ouvrir » le logiciel IDENTITY en écrivant C:\IDENTITY ou cd C:\IDENTITY si vous avez enregistré le logiciel sur le C:\. Vous devriez voir apparaître une ligne C:\IDENTITY> confirmant que le programme IDENTITY est maintenant ouvert.
4. Tapez maintenant le nom de votre fichier à analyser à la suite de C:\IDENTITY> de la façon suivante : identity filename.csv. Le fichier doit être en format .csv, enregistré dans le répertoire IDENTITY que vous avez créé sur le C:\. La base de données doit être formatée ainsi :

No	Cultivar	Locus1	Locus1	Locus2	Locus2
1	Grüner Veltliner	228	230	134	154
2	Cultivar 101	224	224	122	137
2	144			138	138

Il y a donc deux colonnes pour chaque microsatellite (locus), la première colonne d'un locus contenant la première section de la séquence génétique et la deuxième colonne, la deuxième séquence génétique du locus. Les génotypes incomplets à un locus peuvent être inclus dans l'analyse. Il faut alors laisser en blanc les cases correspondant aux deux séquences du locus. C'est à l'utilisateur que revient la responsabilité de déterminer le nombre minimal de loci à utiliser pour la détermination des paires d'individus semblables. Par exemple, les individus ont été identifiés à l'aide de cinq marqueurs microsatellites. Toutefois, le génotypage n'a pas bien fonctionné et, pour plusieurs individus, des marqueurs microsatellites n'ont pu être amplifiés. Il est possible de garder seulement les génotypes complets sur au moins trois loci et il revient à l'utilisateur d'enlever manuellement les individus qui ne répondent pas à ce critère.

5. Une fois le nom du fichier à analyser indiqué, appuyez sur *Entrée*. Une série de données et de commandes apparaîtra rapidement à l'écran DOS. Lorsque l'analyse est terminée, une ligne *Program finished* apparaîtra. La fenêtre DOS peut alors être fermée.
6. Les résultats de l'analyse sont enregistrés directement dans le répertoire IDENTITY sur le C:\. Une description du contenu de chaque fichier de sortie du logiciel IDENTITY est disponible dans le manuel d'instructions PDF. Dans le cas de l'identification de paires d'individus « identiques », c'est le fichier namefile.SYN qui nous intéresse. Ouvrez ce fichier avec le bloc-notes ou tout autre éditeur de fichier texte. Ce fichier présente une liste, en trois colonnes, des paires d'individus semblables. Il est plus pratique de sélectionner tout ce qu'il y a dans la liste et de le copier dans un fichier Excel. Supprimez la colonne du centre, supprimez tous les signes négatifs de la première colonne et vous obtiendrez ainsi une liste des paires d'individus semblables. Les individus portant des astérisques indiquent que le génotype pour ces individus était incomplet.
7. Cette liste peut être épurée. Par exemple, si vous connaissez le sexe des individus et que vous avez un bon niveau de confiance dans la méthode de sexage, il peut être intéressant de comparer le sexe des individus identifiés comme identiques. Ainsi, toutes les paires d'individus dont le sexe diffère peuvent être éliminées. Aussi, certaines répétitions peuvent être éliminées. Par exemple, nous avons obtenu les paires suivantes :
18 = 45
18 = 56
18 = 78
Il faut premièrement s'assurer que 45 est aussi égal à 56 et à 78 et que 56 est aussi égal à 78. Si c'est le cas, les lignes 45 = 56, 45 = 78 et 56 = 78 peuvent être supprimées, puisque de toute façon, les individus 45, 56 et 78 seront remplacés par 18 dans la base de données.
8. Une fois que la liste de paires d'individus « identiques » a été épurée, il faut maintenant remplacer le nom de l'un des individus de la paire par le nom de l'autre. La substitution doit se faire manuellement.
9. Une fois les substitutions de noms terminées, vous pouvez reconstruire l'historique de capture-recapture en considérant que, par exemple, l'individu 18 a été capturé à trois reprises.

Annexe 4. Mode d'emploi pour le logiciel CloseTest.

Le logiciel CloseTest permet de vérifier le degré de fermeture démographique d'une population selon le test proposé par Stanley et Burnham (1999). Le logiciel teste l'hypothèse nulle d'une population fermée selon le modèle M_t (*time specific*) contre l'hypothèse alternative d'une population ouverte selon le modèle de Jolly-Seber. La nature de la violation de la fermeture peut être distinguée (addition ou soustraction d'individus) et les sessions d'échantillonnage pour lesquelles la supposition de fermeture était respectée ou non sont identifiées.

1. Le programme est disponible gratuitement au <http://www.fort.usgs.gov/Products/Software/cloctest/>. Cliquez sur `ctinstall.exe`. Le programme fonctionne sous Windows 2000 ou XP seulement.
2. Ouvrez CloseTest. La fenêtre d'accueil apparaît. Cliquez sur *Main menu* pour poursuivre. Une petite fenêtre comprenant cinq options apparaîtra. L'option *Open* permet de sélectionner la base de données à analyser. L'option *View results* permet de consulter les résultats sauvegardés des tests précédents. L'option *View PDF* permet de consulter l'article de Stanley et Burnham (1999) qui explique les détails du test réalisé par le logiciel. L'option *Help* donne accès à la rubrique d'aide. L'option *Exit* permet de sortir du programme.
3. Cliquez sur *Open* et sélectionnez la base de données à analyser. Les données doivent être dans le format suivant :

0,1,1,0,1
 0,1,0,1,1

 où chaque ligne correspond à un ours et le nombre de colonnes correspond au nombre de sessions d'échantillonnage. Le fichier doit être enregistré en format texte (le format .csv permet d'avoir des virgules entre chaque chiffre).
4. Une deuxième fenêtre apparaîtra, contenant les données et une liste d'options dans la barre du haut. Il est possible de corriger des erreurs directement dans cette fenêtre, à condition de cliquer sur *Re-load* lorsque les corrections sont terminées.
5. Cliquez sur *Analyze* pour réaliser le test de fermeture. Une fenêtre, comme celle présentée ci-dessous, apparaîtra avec les résultats du test.

Data Input File= F:\Sabrina\Résultats Sabrina\Résultats 2000\Closetest\Historique de capture L2000 from identity.csv

N hat= 4191

M_{t+1} = 106

Occasions= 7

CH Data Format= List-directed input

Stanley et Burnham Closure Test (Low p-values suggest population not closed):

Chi-square statistic= 3.99723

df= 5.

p-value= 0.54981

Otis et al. (1978) Closure Test (Low p-values suggest population not closed):

z-value= -999.00000

p-value= 2.00000

Component Statistics of Stanley et Burnham Closure Test

Component	Chi-square	df	p-value
Tests for additions to population (Low p-values suggest there were additions)			
NR vs JS	0.00000	0.	1.00000
M_t vs NM	-1.00000	0.	1.00000
Tests for losses from population (Low p-values suggest there were losses)			
M_t vs NR	3.99723	5.	0.54981
NM vs JS	0.00000	0.	1.00000

Subcomponent Statistics of the NR vs JS Test

(Low p-values on the j-th occasion indicates there were additions to the population between occasions j and j+1)

Occasion	Chi-square	df	p-value
2	Insufficient data for test		
3	Insufficient data for test		
4	Insufficient data for test		
5	Insufficient data for test		
6	Insufficient data for test		

Subcomponent Statistics of the NM vs JS Test

(Low p-values on the j-th occasion indicates there were losses from the population between occasions j-1 and j)

Occasion	Chi-square	df	p-value
2	Insufficient data for test		
3	Insufficient data for test		
4	Insufficient data for test		
5	Insufficient data for test		
6	Insufficient data for test		

- La sortie des résultats donne plusieurs informations intéressantes sur la base de données. Si on s'intéresse seulement au résultat du test de fermeture, il faut aller directement aux sections *Component Statistics of Stanley et Burnham Closure Test*, *Subcomponent Statistics of the NR vs JS Test* et *Subcomponent Statistics of the NM vs JS Test*. Ce sont les résultats

de ces tests qui nous indiquent si la prémisse de fermeture est respectée. Si ce n'est pas le cas, les résultats indiquent s'il y a eu addition ou retrait d'individus et dans quelles occasions de capture la supposition n'est pas respectée. De petites valeurs de p indiquent que la supposition de fermeture n'a pas été respectée.

7. L'exemple ci-dessus montre un cas où la base de données ne contenait pas suffisamment de données pour que le test de fermeture se réalise correctement. Les valeurs de $p = 1.00000$ suggèrent que le *Component Statistics of Stanley et Burnham Closure Test* ne s'est pas bien réalisé et les *Subcomponent Statistics of the NR vs JS Test* et *Subcomponent Statistics of the NM vs JS Test* indiquent qu'il n'y a pas suffisamment de données pour faire le test. Dans ce cas, il est impossible de déterminer si la population à l'étude était fermée démographiquement ou non. Vous pouvez enregistrer la fenêtre de sortie de l'analyse en cliquant sur *File/Save as*. Le fichier sera sauvegardé en format texte.

Annexe 5. Mode d'emploi pour MARK.

Le logiciel MARK permet d'analyser les historiques de capture-recapture et d'estimer le taux de survie et la taille d'une population (N) ouverte ou fermée. Il prend en charge plusieurs formats de données, comme les événements de recapture provenant d'animaux retrouvés morts ou chassés, ou d'animaux réobservés, mais pas recapturés. Pour une revue complète des possibilités d'analyses de MARK et de son fonctionnement, référez-vous au manuel électronique disponible au

<http://warnercnr.colostate.edu/~gwhite/mark/mark.htm#Documentation> en cliquant sur le lien hypertexte [Program MARK A Gentle Introduction](#).

1. Le programme MARK est disponible gratuitement à l'adresse suivante : <http://warnercnr.colostate.edu/~gwhite/mark/mark.htm#Downloading>.

Pour télécharger MARK, cliquez sur *setup.exe*. Vous pouvez enregistrer le programme dans le répertoire C:\Program file.

2. Une fois le programme installé, ouvrez MARK. La fenêtre d'accueil apparaîtra. Cliquez sur *File/New File* ou sur le bouton représentant une feuille avec un tableau superposé juste en dessous de *File*. La fenêtre de travail apparaîtra.

3. Sélectionnez le type de données dans la liste à gauche de l'écran. Dans le cas des données de CMR par génotypage des poils de l'ours noir, le type de données à choisir est *Closed Captures*. Une fois l'option *Closed Captures* sélectionnée, il y aura une fenêtre listant les différentes sous-options possibles. Choisissez l'option qui convient et cliquez sur OK. Dans le cas qui nous intéresse, la sous-option à choisir est *Full Closed Captures with Heterogeneity*.

4. Dans la section de droite de la fenêtre de travail, donnez un titre à votre jeu de données. Sélectionnez le fichier de travail avec le bouton *Click to Select File*. Le nom du fichier apparaîtra dans la case en dessous et automatiquement dans la case *Results File Name*. Il est possible de visualiser la table de données en cliquant sur *View File*, qui deviendra active dès que le fichier de travail sera sélectionné. Le format du fichier doit être le suivant :

Ex. : 11110 12;
10110 8; etc.

La totalité de l'historique de capture doit être contenue dans une seule colonne, où « 1 » représente une capture et « 0 » signifie aucune capture. Le dernier chiffre représente le nombre de fois où le même historique de capture apparaît dans la population. Il est aussi possible de travailler avec un format d'historique de capture où chaque ligne représente un individu et donc que le dernier chiffre avant le point-virgule est un « 1 ». Le fichier doit être en format « namefile.inp ». Pour créer un fichier « .inp », il faut utiliser

un éditeur de texte comme le bloc-notes ou EditPad Lite. Vous pouvez d'abord enregistrer votre historique en format texte avec tabulation (en s'assurant que les valeurs de l'historique de capture apparaissent dans la même colonne au départ pour qu'ils ne se retrouvent pas séparés par des tabulations). Ouvrez le fichier texte avec un éditeur de texte, puis enregistrez-le en ajoutant .inp à la fin du nom de votre fichier.

5. Une fois que la table de données a été sélectionnée, vérifiez que les options dans le coin en bas à droite de l'écran de travail de MARK sont adéquates. Ces options permettent de déterminer le nombre d'occasions de capture, d'ajouter des attributs de groupe et de les nommer, de déterminer l'intervalle entre chaque occasion de capture si ce dernier n'est pas constant et de déterminer le nombre de « mixture ». Il est à noter qu'avec le type de données *Full Closed Captures with Heterogeneity*, il est impossible de spécifier les intervalles de temps entre les occasions de capture.
6. Une fois toutes les options de base déterminées, appuyez sur *OK* dans le coin en bas à droite de l'écran. Un message apparaîtra spécifiant qu'un fichier a été créé avec les résultats. Appuyez sur *Accept*.
7. Une fenêtre rectangulaire nommée *Probability of Mixture* (π) devrait apparaître avec une seule case portant un « 1 ». Cette fenêtre montre le nombre (une case) de paramètres à estimer et le nom (1) du paramètre p_i à estimer. Vous pouvez ouvrir les fenêtres des autres paramètres p (probabilité de capture), c (probabilité de recapture) et N (abondance) en sélectionnant l'option *PIM* de la barre d'outils et en choisissant *Open Parameter Index Matrix*. Une fenêtre apparaîtra vous permettant de choisir les paramètres pour lesquels vous voulez visualiser les fenêtres.
8. La base de la création de modèles dans MARK est la modification de la *PIM Chart* et de la *design matrix*. Il est souhaitable de travailler exclusivement avec la *design matrix* puisque qu'elle permet de faire tous les modèles disponibles pour les *Full Closed Captures with Heterogeneity*, alors que la *PIM Chart* permet seulement de faire les modèles les plus simples. Pour mieux comprendre ce qui constitue une *PIM Chart* et une *design matrix*, référez-vous aux chapitres 6 et 14 du manuel électronique de MARK.
9. Pour débiter la construction d'une *design matrix*, cliquez sur *Design/Reduced*. Sachez qu'il est impossible de construire une *Full matrix* avec le format de données *Full Closed Captures with Heterogeneity*. Une petite fenêtre *Input request* apparaîtra, vous demandant le nombre de colonnes nécessaires dans votre matrice. Le logiciel MARK met, par défaut, le nombre de colonnes nécessaires à la construction de la matrice la plus complexe. Appuyez sur *OK*.

10. Une matrice vide apparaîtra. Cliquez sur *Appearance* pour nommer les colonnes de la matrice. Cette étape est très importante puisque les étapes suivantes consisteront à supprimer des colonnes et que pour cela, il est essentiel de les nommer. Si, par exemple, vous avez quatre occasions de capture, vous aurez la séquence de colonnes suivante : PI, INTERCEPT, E, H, T1, T2, T3, HT1, HT2, HT3, ET2, ET3, HET1, HET2, HET3, N.

Où :

- PI correspond au paramètre π et INTERCEPT à l'ordonnée à l'origine du modèle;
- E réfère à *Encounter group*, une variable qui permet de modéliser la réponse comportementale des individus entre les événements de capture (*encounter*) et qui n'apparaît seulement que dans les modèles avec variation interindividuelle;
- H réfère à la composante de variation dans la probabilité de capture dans le temps;
- T réfère au nombre de sessions de capture. Si on a quatre sessions, on a T1 à T3;
- Les autres colonnes représentent des interactions entre les différentes variables;
- N représente l'estimation d'abondance.

11. Une fois les colonnes nommées, appuyez sur OK. Il faut maintenant construire manuellement la matrice. Utilisez la page 26 du chapitre 14 du manuel électronique de MARK comme référence. Remplissez les cases avec des 1 en calquant le modèle de matrice de la page 26. Ajustez ce modèle en fonction du nombre de sessions de capture. Remarquez que les variables d'interaction entre E et T débutent à ET2 et non à ET1 et que les variables d'interaction entre H, E et T débutent quant à elles à 1. Cette matrice constitue la matrice la plus complète, toutefois il est impossible d'estimer N avec cette matrice. En effet, une contrainte doit être imposée au paramètre p (probabilité de capture) pour la dernière occasion de capture (voir la section 14.3.1 du manuel pour plus de détails sur la contrainte sur le dernier p). Pour mettre une contrainte sur le dernier p dans le cadre des modèles *Full Closed Captures with Heterogeneity*, il faut supprimer l'interaction entre E et T et entre H et T. Il faut donc supprimer les colonnes ET1 à ETi, HET1 à HETi et HT1 à HTi en utilisant la fonction *DelCol/One column* ou *DelCol/Multiple columns*.

12. La matrice nouvellement formée constitue le modèle le plus complexe, M_{tth} , qu'il est possible de construire avec le type de données *Full Closed Captures with Heterogeneity*. Si les modèles ne contiennent pas la variable h , il faut passer au type de données *Closed Capture with Heterogeneity*; dans le cadre des modèles simples M_t , M_b et M_o , il faut passer au type de données *Closed Capture*. Pour changer le type de données, cliquez sur *PIM, Change data type*. La liste de tous les types de données disponibles apparaîtra. La liste complète des modèles qu'il est possible de réaliser avec ce type de données est présentée ci-dessous :

Nom du modèle	Notation
M_{tbh}	$\{N, p_i, p(t), c(t)\}$
M_{tb}	$\{N, p(t), c(t)\}$
M_{th}	$\{N, p_i, p(t)=c(t)\}$
M_{bh}	$\{N, p_i, p(.), c(.)\}$
M_t	$\{N, p(t)=c(t)\}$
M_b	$\{N, p(.), c(.)\}$
M_h	$\{N, p_i, p(.)=c(.)\}$
M_o	$\{N, p(.)=c(.)\}$

Où p correspond à la probabilité de capture, c à la probabilité de recapture, p_i à la présence de « mixture group » et donc à la présence de variabilité dans p et c , et N à l'abondance. La présence de t entre parenthèses indique la présence d'un effet temporel dans les probabilités de capture et de recapture, alors que la présence d'un point réfère à la constance des probabilités dans le temps. Le symbole $=$ entre p et c signifie que les probabilités sont égales, donc qu'il n'y a pas de réponse comportementale de la probabilité de recapture après une première capture.

13. Pour faire exécuter le modèle M_{tbh} que vous venez de construire, cliquez sur *Run/Run*. Une fenêtre *Setup numerical estimation run* apparaîtra, vous permettant de sélectionner les options d'estimation que vous désirez. Commencez par nommer le modèle que vous venez de construire dans la case *Model name*. Le nom n'influence en rien la composition du modèle. Optez pour des noms explicites. Par exemple, le modèle M_{tbh} pour $\{M_{tbh} : N, p_i, p(t), c(t)\}$. Le bloc en bas à gauche de la fenêtre vous permet de choisir la *Link function* que vous désirez. Par défaut, MARK met *Logit* comme fonction pour des *Full Closed Captures with Heterogeneity*, et le manuel suggère de conserver cette fonction par défaut. Même chose pour le bloc *Var. estimation* pour lequel l'option par défaut est *2ndpart* que l'on suggère de conserver par défaut. Le bloc de droite offre d'autres options pour l'estimation. Pour les analyses réalisées dans le présent rapport, aucune autre option n'a été utilisée. Une fois que les options que vous désirez sont déterminées, cliquez sur *OK to run*.
14. Un message apparaîtra, vous demandant si vous voulez ajouter les résultats de ce modèle dans la base de données. Cliquez sur *Yes*. Si vous avez utilisé le format d'historique de capture avec une ligne pour chaque ours se terminant par un 1, le message vous indiquera aussi un message d'avertissement soulignant que plus d'une ligne dans la base de données sont identiques. Ce message d'avertissement est normal, ne vous en préoccupez pas. Cependant, il est possible qu'un autre message d'avertissement apparaisse : ****Warning** error 2 from VA09AD optimization routine.**

Ce message indique qu'un problème est survenu lors de l'optimisation des valeurs des paramètres du modèle (Gary White, Université du Colorado, commun. pers.) qui peut être due à une surparamétrisation du modèle ou une taille d'échantillon trop faible.

15. Une fois que vous aurez ajouté les résultats du modèle dans la base de données (*Result Browser*) vous pourrez visualiser différentes statistiques du modèle, comme dans l'exemple ci-dessous :

Model	AICc	Delta AICc	AICc Weights	Model Likelihood	Num. Par	Deviance
{M _{tbh} , N, p _i , p(t), c(t)}	-2 879,386	0,1032	0,37248	0,9497	11	318 159

16. Pour construire les autres modèles, vous pouvez récupérer la matrice la plus complète (M_{tbh}) ou celle qui est la plus près du prochain modèle que vous voulez construire en cliquant d'abord sur le modèle pour lequel vous voulez reconstituer la matrice (le modèle deviendra bleu) et en cliquant ensuite sur le bouton composé de quatre carrés noirs, au-dessus de la liste de modèles. La matrice réapparaîtra. Il suffit alors de supprimer les colonnes nécessaires pour créer un nouveau modèle. Il est à noter que lorsque vous changez de type de données (*Closed Capture with Heterogeneity* ou *Closed Capture*) il est possible que vous ayez à reconstruire une matrice avec le nouveau format (modèle M_{tb}) ou que vous ayez à construire la matrice à l'aide de la *PIM Chart* (modèle M_h , M_t , M_b , M_o). Pour ouvrir une *PIM Chart*, allez dans l'onglet *PIM*, cliquez sur *Parameter index chart*. Un graphique comportant les paramètres en ordonnée et des carrés bleus devrait apparaître. À partir de l'onglet *Initial*, vous pouvez choisir les options possibles pour construire les modèles M_h , M_t , M_b et M_o , en vous basant sur les informations fournies dans le tableau du point 13 ci-haut.

Modèle	Colonnes à conserver dans la matrice
M_{th}	PI, INTERCEPT, H, T1 à Ti, N
M_{tb}	INTERCEPT, E, T1 à Ti, N
M_{bh}	PI, INTERCEPT, E, H, N

17. Lorsque vous aurez fait tous les modèles que vous désirez, votre *Results Browser* devrait ressembler à une table comme celle-ci :

Model	AICc	Delta AICc	AICc Weights	Model Likelihood	Num. Par	Deviance
{Mtb, N, p(t), c(t)}	-2879,4892	0	0,39220	1	11	31,7128
{Mtbh, N, pi, p(t), c(t)}	-2879,3860	0,1032	0,37248	0,9497	11	31,8159
{Mth, N, pi, p(t)=c(t)}	-2878,4675	1,0217	0,23532	0,6	10	34,7528
{Mt, N, p(t)=c(t)}	-2847,7420	31,7472	0	0	6	73,5348
{Mbh, N, pi, p(.), c(.)}	-2802,8508	76,6384	0	0	5	120,4359
{Mb, N, p(.), c(.)}	-2795,4059	84,0833	0	0	3	131,8958
{Mh, N, pi, p(.)=c(.)}	-2728,2303	151,2589	0	0	4	197,0647
{M0, N, p(.)=c(.)}	-2722,7711	156,7181	0	0	2	206,5355

Vous pouvez copier et enregistrer ce tableau en cliquant sur le menu *Output/Table of model results/Copy to Excel*. Vous pouvez aussi consulter les sorties pour chaque modèle individuellement en cliquant sur le bouton à droite de la corbeille dans le haut de la fenêtre de *Result Browser*. Vous pouvez voir les estimations pour chacun des paramètres du modèle en cliquant sur le quatrième bouton à partir de la gauche, et l'estimation d'abondance (N-hat) en cliquant sur le cinquième bouton (bouton bleu) à partir de la gauche.

18. Un des avantages de MARK est la possibilité de réaliser une estimation de N avec une moyenne pondérée selon le poids de chacun des modèles. Pour ce faire, cliquez sur *Output/Model averaging/Real*. Une fenêtre *Parameter selection for model averaging* apparaîtra. Dans la liste déroulante de *Select PIM*, choisissez le paramètre pour lequel vous voulez avoir une estimation. Dans le coin en haut à droite de la fenêtre, cliquez sur *Add to list*. Le paramètre apparaîtra dans la grande case blanche *Parameter to model average*. Vous pouvez ajouter le nombre de paramètres que vous voulez. Cliquez ensuite sur *OK*.

19. Le bloc-notes s'ouvrira avec les résultats du « model averaging ». La moyenne pondérée apparaît tout en bas du tableau, avec son erreur standard inconditionnelle (*unconditional SE*). Si, pour certains modèles, l'estimation du paramètre ne peut pas se réaliser correctement, vous risquez de voir apparaître, par exemple pour N, une estimation égale à M(t+1) et une SE égale à 0. Vous devriez alors avoir dans la sortie des résultats du « model averaging » un message *Invalid parameter estimate with zero SE?* à droite. La valeur de M(t+1) peut être visualisée en cliquant sur le deuxième bouton à partir de la gauche dans le *Result Browser*. Ces modèles devraient être éliminés de la liste de modèles candidats pour le « model averaging ». Pour supprimer ces modèles, cliquez d'abord sur le modèle et cliquez ensuite sur le bouton corbeille en haut à gauche du *Result Browser*. Refaites ensuite le « model averaging ».

20. Pour calculer un intervalle de confiance pour une estimation issue du « model averaging », le manuel de MARK suggère une méthode de calcul qui utilise l'erreur standard inconditionnelle fournie dans la sortie de résultats du « model averaging ». Ce calcul est montré à la page 31 du chapitre 14.

Annexe 6. Mode d'emploi pour DENSITY 4.4.

Le logiciel DENSITY 4.4 permet d'estimer la taille et la densité d'une population en tenant compte de la disposition spatiale des stations d'échantillonnage. Dans ce rapport, nous avons utilisé DENSITY pour obtenir la superficie effective de l'aire d'étude pour calculer une densité de population à partir d'une estimation de taille de population issue du « model averaging » du programme MARK.

1. Le logiciel DENSITY 4.4 est disponible gratuitement sur le site <http://www.otago.ac.nz/density/>. Cliquez sur *Download 4.4* pour le télécharger. Vous devez remplir la liste d'informations requises pour télécharger le logiciel.
2. Ouvrez le logiciel DENSITY 4.4. Une fenêtre d'accueil s'ouvrira. Cliquez sur *Continue*.
3. Comme mentionné, le logiciel calculera une densité de population en tenant compte de l'arrangement spatial des stations d'échantillonnage. Vous devez donc créer un fichier contenant les coordonnées géographiques des stations. Le fichier doit contenir trois colonnes, à savoir TRAP ID, X COORD, Y COORD (en UTM), en respectant cet ordre. Le fichier doit être en format texte et ne doit pas contenir les titres des colonnes. Cliquez sur le bouton contenant l'icône d'un fichier dans la case TRAP LAYOUT. Vous pourrez sélectionner le fichier de coordonnées de vos stations. Dans la liste déroulante *Type*, cliquez sur *Proximity* si vos stations permettent d'échantillonner plus d'un individu par occasion de capture et d'échantillonner un même individu à plusieurs stations dans une même session d'échantillonnage (c'est le cas des pièges à poils utilisés avec l'ours). Dans la case *Buffer (m)*, tapez le nombre de mètres de la zone tampon que vous voulez voir apparaître sur la carte autour des stations. Vous pourrez faire varier la taille de cette zone plus tard si vous le désirez.
4. Dans la case CAPTURE DATA, vous devez sélectionner le fichier contenant l'historique de capture. Cette base de données doit avoir le format SESSION ID, BEAR ID, OCCASION ID, TRAP ID (les titres ne sont qu'à titre indicatif et ne doivent pas apparaître dans le fichier texte). La différence entre les sessions et les occasions de capture est que les sessions sont les périodes larges de capture (p. ex., l'été 2001 représente une seule session), alors que les occasions sont les différentes périodes de visite des stations (p. ex., si les stations ont été visitées durant cinq semaines consécutives à l'été 2001, il y aura cinq occasions de capture). Chaque ligne représente une capture et non un ours. Il peut donc y avoir plusieurs lignes avec le même BEAR ID, mais à des stations ou des occasions de capture différentes. Une fois que le fichier a été chargé, vous devez spécifier le type de données que vous avez dans la liste déroulante *Type*. Si vous utilisez une colonne TRAP ID dans votre fichier d'historique de capture, alors vous devez

spécifier TRAP ID dans *Type*. Spécifiez XY si, dans cette base, vous avez utilisé les positions géographiques des stations plutôt que simplement leur identité.

5. Cliquez sur le bouton *Option* dans le coin supérieur gauche de l'interface de travail. Dans l'onglet *Input*, sélectionnez le nombre d'occasions de capture. Cliquez aussi sur la case *Allow recapture within an occasion* si vous avez des individus qui ont été capturés plus d'une fois durant une occasion. Dans l'onglet *Output*, cliquez la *sq km* pour obtenir les estimations de l'aire effective et de la densité en km².
6. Cliquez sur le bouton *Read data* dans le coin supérieur gauche de l'interface de travail. Il est possible que des messages apparaissent, vous indiquant qu'il existe des doubles dans la liste de noms des stations. DENSITY s'occupe de les enlever si vous appuyez sur OK. Si des stations se trouvent dans l'historique de capture, mais qu'elles sont absentes de la liste de localisations, un message apparaîtra. Ce sera à vous de corriger cette erreur, soit en ajoutant la position de la station dans la liste de localisations si cette dernière est connue, ou en enlevant l'évènement de capture correspondant à cette station si sa position est inconnue. Si tout est en ordre, un bloc contenant plusieurs onglets devrait apparaître à droite dans l'interface. Une carte devrait aussi apparaître sous le bloc *CAPTURE DATA*, montrant la disposition spatiale des stations. Il est possible d'afficher les captures en appuyant sur l'option *Capture* en haut de l'écran. Il est aussi possible d'afficher le nom des stations et les déplacements en cliquant à droite sur la carte et en sélectionnant l'option dans la liste.
7. Pour obtenir **l'aire d'étude effective**, cliquez sur l'onglet *ETA density* dans la section de droite du module. À partir de cet onglet, il est possible d'obtenir l'aire effective d'étude à l'aide de quatre calculateurs (strip methods) : *Manual*, *ARL/2*, *MMDM*, et *MMDM/2*. Pour l'option *Manual*, vous devez entrer la largeur de la bande entourant les stations que vous désirez. Par exemple, dans le cas de l'étude de 2000-2003 en Abitibi-Témiscamingue, nous avons utilisé une bande de 5 046 m, ce qui correspond au diamètre d'un domaine vital circulaire de 20 km². Il est possible que les données soient insuffisantes pour calculer *ARL/2* et même *MMDM* et *MMDM/2*. Dans ce cas, il sera impossible de déterminer l'aire effective d'étude avec ces méthodes. Par défaut, DENSITY met les estimations d'aire effective d'étude en hectares (ha). Si vous voulez que DENSITY estime ces aires en km², cliquez sur le bouton *OPTION* en haut à gauche de l'écran et cliquez la case km².
8. Pour obtenir une estimation de **l'abondance**, cliquez sur l'onglet *Population*. Cet onglet offre différents estimateurs permettant de calculer l'abondance. Choisissez le modèle désiré et appuyez sur le bouton *GO*, au centre en haut de l'interface. L'estimation et son intervalle de confiance apparaîtront à côté de Population size $N\text{-hat} \pm SE$ (IC 95%).

Cliquez la case *Append output* et cliquez sur *RUN*. Vous pouvez visualiser les résultats en cliquant sur les cases *View Output* dans l'onglet ou dans le haut de l'écran. Vous pouvez aussi consulter le bloc *View log* pour vérifier si le test s'est bien réalisé. Ces résultats sont enregistrés dans des fichiers textes nommés *density4.out* et *.log*, que vous devez renommer et enregistrer si vous ne voulez pas les effacer lors des prochaines analyses. Dans le bloc *OPTION* en haut à gauche de l'écran, vous pouvez spécifier l'endroit où vous voulez que ces fichiers soient enregistrés.

9. Pour réaliser une **ML SECR** (*maximum likelihood spatial explicit capture-recapture*), cliquez sur la petite boîte à côté de ML SECR dans le bloc *ANALYSIS GROUPS*. Cliquez sur la boîte d'options juste à gauche du nom ML SECR et choisissez les options qui vous intéressent. Les différents modèles testés sont sélectionnés ici en cliquant dans les petites cases circulaires correspondant aux paramètres que vous voulez ajouter dans le modèle. Une fois que toutes les options sont choisies, appuyez sur *GO* dans le haut de l'écran. Selon le nombre de paramètres inclus dans le modèle, il est possible que le logiciel prenne un certain temps à résoudre le modèle. La réussite du test et les résultats apparaîtront dans les fichiers *Density.log* et *Density.out*. Dans le *Density.out*, vous pouvez retirer les informations suivantes : *MLDensity* (densité, en individu par km²), *MLAIC* et *MLAICc* (les valeurs d'AIC et d'AICc), *ML.SE* (l'erreur type), *ML.LC* (la limite inférieure de l'intervalle de confiance) et *ML.UC* (la limite supérieure de l'intervalle de confiance). Pour plus d'informations, voir l'article d'Efford (2009) et la section ML SECR du site web <http://www.otago.ac.nz/density/>.

Annexe 7. Mode d'emploi pour le paquet SECR dans le logiciel R 2.12.2.

Le paquet *secr* permet, à l'aide de l'environnement de travail R, d'estimer la densité avec la méthode de la maximisation de la vraisemblance pour des données de capture-recapture spatialement explicites (ML SECR). L'utilisation du logiciel R pour réaliser une ML SECR offre plusieurs avantages par rapport au logiciel DENSITY 4.4. Entre autres, le test de fermeture des populations peut être fait avec ce logiciel et il est possible d'estimer la densité avec des tests non spatiaux pour des populations fermées. Il est aussi possible de représenter graphiquement la fonction de détection g des individus et de cartographier le centre estimé des domaines vitaux des individus et le contour de la probabilité de détection. Pour une liste exhaustive des différences entre le logiciel DENSITY 4.4 et le paquet *secr* du logiciel R, voir Efford (2011).

1. Il est possible de télécharger gratuitement le logiciel R sur le site suivant :

<http://www.t-project.org/index.html>.

2. Ouvrez le logiciel R. Une fois ouvert, vous devez spécifier le répertoire de travail dans lequel les bases de données sont situées (*work directory*). Il est possible de choisir ce répertoire de travail en cliquant sur l'onglet *Files*, puis sur l'option *Change directory*. Une fenêtre apparaîtra vous permettant de choisir l'emplacement du répertoire de travail.

3. Chargez le paquet *abind*. Pour se faire, cliquez d'abord sur l'onglet *Packages* et choisissez l'option *Install Package*. Cliquez ensuite sur *Load Package*, une liste des paquets disponibles apparaîtra. Cliquez sur le paquet *abind*. Le logiciel R le chargera automatiquement. L'opération prendra quelques secondes et vous verrez apparaître la commande attestant le chargement du paquet dans l'environnement de travail. Il est possible que le logiciel vous demande de choisir un miroir pour le chargement (*mirror* ou *CRAN mirror*). Choisissez le miroir correspondant au Canada et Québec (1 ou 2, aucune importance).

4. Chargez ensuite la librairie *secr*. Pour ce faire, tapez *library(secr)* dans la fenêtre de travail. Le symbole « > » apparaîtra lorsque la librairie sera chargée. Il est possible que vous deviez charger le paquet *secr* avant. Cette étape permet en fait d'avoir accès à toutes les fonctions et aux sections d'aide relatives à la méthode de maximisation de la vraisemblance pour des captures-recaptures spatialement explicites (ML SECR).

5. Avant de commencer les analyses, vous devez d'abord avoir en mains un historique de capture en format .txt avec le format suivant : SESSION ID, BEAR ID, OCCASION ID, TRAP ID (les titres ne sont qu'à titre indicatif et ne doivent pas apparaître dans le fichier texte, mais les colonnes doivent apparaître dans cet ordre). Vous devez aussi avoir une base de données contenant les coordonnées géographiques des stations d'échantillonnage en format .txt : TRAP ID, X COORD, Y COORD (en UTM, l'ordre des colonnes est important, mais elles ne doivent pas apparaître dans le fichier texte). Il faut utiliser le même format que pour le logiciel DENSITY. Placez ces fichiers dans le répertoire de travail que vous avez précédemment spécifié.

6. Vous devez maintenant créer un historique de capture contenant à la fois l'historique de capture et l'identité de la station où la capture s'est effectuée. Le logiciel R est capable de créer ce format d'historique de capture avec la fonction *read.caphist*. Choisissez d'abord un nom pour l'historique de capture (dans l'exemple qui suit, le nom choisi pour l'historique de capture est *caph*). Ensuite, tapez la séquence suivante :

```
caph<-read.caphist("histR2001.txt","trap2001.txt",detector="proximity",fmt="trapID")
```

où « histR2001.txt » est l'historique de capture construit au point 5 (voir ci-haut) et « trap2001.txt » est la base de données de coordonnées géographiques des stations. L'argument *detector* représente le type de station que vous avez. Dans le cadre d'une étude de CMR avec génotypage des poils, les stations sont considérées comme des détecteurs « proximity » s'il n'y a pas d'évènement de « double détection » pour un même individu dans la même session d'échantillonnage. Si, au contraire, vous avez des évènements de « multiples détections », vous devez choisir le type de détecteur « count ». Vous pouvez d'abord utiliser le type de détecteur « proximity » sans vérifier la présence d'évènements de « multiples détections ». Si vous avez des évènements de « multiples détections », le logiciel vous le fera remarquer en affichant le message d'erreur suivant lors de la commande *read.caphist* :

Session 1

More than one detection per detector per occasion at proximity detector(s)

L'argument *fmt* spécifie le format de l'historique de capture. Dans cet exemple, l'identité des stations plutôt que leurs coordonnées géographiques est utilisée dans l'historique de capture. Pour plus de détails sur les options de la fonction *read.caphist*, tapez *?read.caphist*. Si la séquence a bien été entrée, vous devriez voir le message suivant apparaître :

>no errors found :-)

7. Vous devez maintenant exécuter les modèles de populations fermées que vous désirez. Vous ferez exécuter les modèles à l'aide de la fonction *secur.fit*. Les arguments essentiels à la séquence sont le nom de la base de données *caphist* que vous avez créée auparavant, le modèle (les paramètres) que vous voulez tester, la largeur de la zone tampon entourant les stations (*buffer*). Cette zone tampon est en fait la distance sur laquelle les fonctions de détection (appelées *g0*) et la fonction de décroissance de *g0* (appelée *sigma*) s'étendent. Le logiciel R peut suggérer une superficie de zone tampon selon la base de données à traiter avec la fonction *suggest.buffer*. Pour ce faire, tapez la séquence suivante :

```
suggest.buffer(caph, detectfn=NULL, detectpar=NULL)
```

où *caph* représente le nom de la base de données *caphist* créée antérieurement et où *detectfn* précise la forme de la fonction de détection *g0*. Dans cet exemple, *g0* est considérée comme une fonction demi-normale qui se traduit par l'argument *NULL*. L'argument *detectpar* doit représenter la liste des paramètres à évaluer avec la ML SECR

(g_0 et σ). Pour les analyses incluses dans le présent rapport, l'utilisation de l'argument `NULL` s'est avérée correcte puisque le logiciel R utilise alors des valeurs automatiques. Aussi, par défaut, le logiciel R présume que la distribution du nombre de détections par individu par station suit une loi de Poisson uniforme, ce qui est correct. Nous n'avons donc rien à spécifier en ce qui a trait à cette fonction. Vous devriez obtenir une sortie de résultats de ce genre :

```
> suggest.buffer(capth30, detectfn=NULL, detectpar=NULL)
```

```
[1] 6003
```

```
using automatic 'detectpar' g0 = 0.1402, sigma = 1877
```

La bande entourant les stations pour la suite des analyses sera de 6 003 m.

8. Vous êtes maintenant prêt à exécuter les modèles qui vous désirez. Vous utiliserez la fonction *secr.fit*. Par exemple pour le modèle simple $g(.)s(.)$, vous devez taper :

```
secr00<-secr.fit(capth, model=list(g0~1, sigma~1), buffer=6003, trace=FALSE)
```

où *capth* est l'historique de capture comprenant l'identité des stations qui a été créé précédemment, l'argument *model* permet de lister les paramètres à évaluer, l'argument *buffer* permet de spécifier la largeur de la zone tampon entourant les stations suggérée par le logiciel et l'argument *trace* permet de spécifier si nous voulons obtenir les résultats de vraisemblance pour chaque évaluation (`FALSE` les cache, `TRUE` les fournit). Le tableau 1 illustre la syntaxe pour la combinaison des paramètres de temps (t), de réponse comportementale à une première capture (b) et de variation interindividuelle dans la probabilité de capture (h_2).

L'analyse du modèle sera terminée lorsque le symbole « > » sera réapparu. Il est possible que des messages d'avertissement s'affichent. Par exemple :

```
In secr.fit(capth, model = list(g0 ~ 1, sigma ~ h2), buffer = 7263, :  
  predicted relative bias exceeds 0.01 with buffer = 7263
```

Ce message suggère que le logiciel a été incapable de maximiser la vraisemblance avec les paramètres suggérés (ici h_2 en σ) en limitant la taille du biais (biais supérieur à 0,01 ou 1 % de la valeur prédite). Les résultats d'un tel modèle doivent donc être interprétés avec précautions.

Tableau 1. Syntaxe pour les différentes combinaisons des variables t, b et h2 selon les fonctions de détection g0 et sigma, seulement pour l'argument *model* de la fonction *secr.fit*.

Modèle	Syntaxe pour l'argument Model dans R
g(.)s(.)	model=list(g0~1, sigma~1)
g(.)s(h)	model=list(g0~1, sigma~h2)
g(t)s(.)	model=list(g0~t, sigma~1)
g(t)s(h)	model=list(g0~t, sigma~h2)
g(b)s(.)	model=list(g0~b, sigma~1)
g(b)s(h)	model=list(g0~b, sigma~h2)
g(h)s(.)	model=list(g0~h2, sigma~1)
g(h)s(t)	model=list(g0~h2, sigma~h2)
g(tb)s(.)	model=list(g0~t+b, sigma~1)
g(tb)s(h)	model=list(g0~t+b, sigma~h2)
g(th)s(.)	model=list(g0~t+h2, sigma~1)
g(th)s(h)	model=list(g0~t+h2, sigma~h2)
g(bh)s(.)	model=list(g0~b+h2, sigma~1)
g(bh)s(h)	model=list(g0~b+h2, sigma~h2)
g(tbh)s(.)	model=list(g0~t+b+h2, sigma~1)
g(tbh)s(h)	model=list(g0~t+b+h2, sigma~h2)

9. Pour voir les résultats, tapez le nom du modèle que vous aviez choisi, par exemple *secr0h2* :

```
> secr0h2
```

Voici un exemple de ce que vous devriez obtenir comme résultats :

```
secr.fit( capthist = capth, model = list(g0 ~ 1, sigma ~ h2), buffer = 7263, trace = FALSE )
secr 2.2.0, 14:58:41 02 abr 2012
Detector type  proximity
Detector number 271
Average spacing 3046.73 m
x-range      610340 698431 m
y-range      5290998 5375103 m
N animals    : 476
N detections  : 619
N occasions  : 5
Mask area    : 1020118 ha
```

Model : D~1 g0~1 sigma~h2 pmix~h2
 Fixed (real) : none
 Detection fn : halfnormal
 Distribution : poisson
 N parameters : 5
 Log likelihood : -1536.782
 AIC : 3083.565
 AICc : 3083.692

Beta parameters (coefficients)

	beta	SE.beta	lcl	ucl
D	-5.193412	0.14482920	-5.477272	-4.909552
g0	-1.454334	0.11197360	-1.673798	-1.234870
sigma	8.335252	0.06861978	8.200760	8.469744
sigma.h22	-1.656184	0.08513882	-1.823053	-1.489315
pmix.h22	5.616658	0.32929186	4.971258	6.262058

Variance-covariance matrix of beta parameters

	D	g0	sigma	sigma.h22	pmix.h22
D	0.0209754973	-0.004291927	-0.0005943137	-0.007857363	0.012846222
g0	-0.0042919267	0.012538087	-0.0028026522	0.000776626	-0.003561535
sigma	-0.0005943137	-0.002802652	0.0047086744	-0.003403695	0.004798388
sigma.h22	-0.0078573634	0.000776626	-0.0034036947	0.007248618	-0.009021365
pmix.h22	0.0128462224	-0.003561535	0.0047983878	-0.009021365	0.108433128

Fitted (real) parameters evaluated at base levels of covariates

session = 1, h2 = 1

	link	estimate	SE.estimate	lcl	ucl
D	log	5.553025e-03	8.084760e-04	4.180717e-03	7.375789e-03
g0	logit	1.893354e-01	1.718655e-02	1.579184e-01	2.253302e-01
sigma	log	4.168253e+03	2.863616e+02	3.643718e+03	4.768297e+03
pmix	logit	3.623596e-03	NA	NA	NA

session = 1, h2 = 2

	link	estimate	SE.estimate	lcl	ucl
D	log	5.553025e-03	0.000808476	4.180717e-03	7.375789e-03
g0	logit	1.893354e-01	0.017186551	1.579184e-01	2.253302e-01
sigma	log	7.955774e+02	54.656675664	6.954616e+02	9.101054e+02
pmix	logit	9.963764e-01	NA	NA	NA

Bien que la sortie de résultats soit longue, les informations réellement importantes dans l'estimation de la densité d'une population sont l'AICc (ou l'AIC), estimate, SE.estimate, lcl et ucl de la section *Fitted (real) parameters evaluated at base levels of covariates*. « Estimate » est l'estimation de densité en individu par hectare. Pour convertir en nombre

d'individus par 10 km², multipliez l'estimation par 1 000. Dans cet exemple, la densité serait de 0,5530 individu/km². « SE.estimate » représente l'erreur type (standard error) de l'estimation, laquelle doit aussi être multipliée par 1 000 si on désire exprimer la densité en individu/10 km². « Lcl » et « ucl » représentent respectivement les limites inférieure et supérieure de l'intervalle de confiance (qui doivent aussi être multipliées par 1 000). Remarquez qu'il y a deux sections de résultats : session=1, h2=1 et session=1, h2=2. Ceci est dû à la composante de variation utilisée dans ce modèle, composante appelée 2-class finite-mixture, ce qui produit deux tableaux de résultats pour les deux classes de variation. Les résultats de ces deux sections seront inévitablement les mêmes, donc peu importe de quelle section proviennent les résultats.

La fonction *AIC** peut être utilisée pour obtenir les poids d'Akaike au lieu de les calculer. Il est aussi possible de calculer la moyenne pondérée de l'estimation de densité avec la fonction *model.average*. Vous devez toutefois exécuter tous les modèles désirés sans fermer la fenêtre de travail R pour utiliser la fonction *model.average* puisque cette dernière exige que l'on précise tous les noms des modèles que l'on veut inclure dans le calcul de la moyenne pondérée. En fermant la fenêtre de travail R, vous perdrez l'accès à ces modèles. Il en est de même pour la conservation des résultats. Nous vous recommandons le logiciel Tinn-R ou le bloc-notes pour copier/coller les résultats à mesure qu'ils sont obtenus.

Il est possible que vous obteniez une variété de messages d'erreur. Voici les plus fréquents :

```
In secr.fit(capth, model = list(g0 ~ 1, sigma ~ h2), buffer = 7263, : predicted relative bias exceeds 0.01 with buffer = 7263
```

Ce message suggère que le biais de l'estimation de densité avec la largeur de la zone tampon autour des stations est plus élevé que la limite déterminée de 0,01 (1 %). Ainsi, on accorde peu de confiance à l'estimation obtenue. Tous les messages d'erreur suivants, bien que l'erreur provienne de sources différentes, suggèrent la même chose : la maximisation de la vraisemblance n'a pu être réalisée correctement à cause d'un manque de données et/ou d'une surparamétrisation du modèle.

```
In secr.fit(capthT2002, model = list(g0 ~ t + h2, sigma ~ 1), buffer = 23414, : variance calculation failed - try method = 'BFGS'
Error en nlsModel(formula, mf, start, wts) : singular gradient matrix at initial parameter estimates
Error en nls(temp3 ~ (1 - (1 - gr(buff))^(noccasions * k)), start = list(k = 2)) : step factor 0.000488281 reduced below 'minFactor' of 0.000976563.
```

Si ces messages apparaissent, il est possible que vous ne puissiez faire le calcul souhaité. Si le logiciel produit tout de même des résultats, ceux-ci risquent d'être imprécis et biaisés et il est conseillé de ne pas les utiliser.

**Forêts, Faune
et Parcs**

Québec

