



Etude de consistance et applications du modèle Poisson-gamma : modélisation d'une dynamique de recrutement multicentrique

Nathan Minois

► To cite this version:

Nathan Minois. Etude de consistance et applications du modèle Poisson-gamma : modélisation d'une dynamique de recrutement multicentrique. Statistiques [math.ST]. Université Paul Sabatier - Toulouse III, 2016. Français. NNT : 2016TOU30396 . tel-02275009

HAL Id: tel-02275009

<https://theses.hal.science/tel-02275009>

Submitted on 30 Aug 2019

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

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



THÈSE

En vue de l'obtention du

DOCTORAT DE L'UNIVERSITÉ DE TOULOUSE

Délivré par :

Université Toulouse 3 Paul Sabatier (UT3 Paul Sabatier)

Présentée et soutenue par :

Nathan Minois

le lundi 7 novembre 2016

Titre :

Étude de consistance et applications du modèle Poisson-gamma : modélisation
d'une dynamique de recrutement multicentrique

École doctorale et discipline ou spécialité :

ED MITT : Domaine Mathématiques : Mathématiques appliquées

Unité de recherche :

Inserm UMR 1027

Directeur/trice(s) de Thèse :

Professeur Sandrine Andrieu

Maître de conférences Nicolas Savy

Jury :

Professeur Jean-François Dupuy - Rapporteur

Maître de Conférences Nicolas Molinari - Rapporteur

Maître de Conférence Valérie Lauwers-Cances - Examineur

Professeur Aurélien Latouche - Examineur

Remerciements

C'est avec grand plaisir que je tiens à remercier tous ceux qui m'ont aidé dans l'aboutissement de ma thèse.

Je tiens tout d'abord à remercier mes deux encadrants de thèse Sandrine Andrieu et Nicolas Savy pour leur soutien tout au long de la réalisation de ce projet ainsi que pour la confiance qu'ils ont su me porter suite à la réalisation de mon stage de fin d'étude de Master II IMAT pour la poursuite en doctorat. Leurs apports de par leurs compétences professionnelles et humaines m'ont permis de travailler dans un cadre idéal, que ce soit en terme de perspectives de recherche ainsi que les multiples présentations lors de congrès tout au long de la thèse.

Je tiens également à remercier Vladimir Anisimov pour son aide sans faille lors de la réalisation de multiples études de par ses analyses pertinentes ainsi que de son aide rédactionnelle et théorique durant l'écriture d'articles.

Je remercie également mes deux rapporteurs Jean-François Dupuy et Nicolas Molinari qui ont grandement contribué à la version finale de ce mémoire avec malheureusement comme tâche principale et ingrate la corrections de multiples fautes d'orthographe (qui doivent encore être nombreuses). Je remercie également Valérie Lauwers-Cances pour avoir accepté de faire partie du jury de thèse ainsi que pour son aide lors de la réalisation de multiples études et articles tout au long du doctorat, et également Aurélien Latouche qui bien que très pris sera présent au sein du jury par visioconférence lors de la soutenance.

Je tiens de plus à remercier l'UMR 1027 de l'Inserm et plus particulièrement l'équipe 1 pour son accueil et son aide précieuse tout d'abord en tant que stagiaire puis en tant que doctorant. Un remerciement particulier pour Gregory Guernec, Gauthier CHASSANG, Michael Mounie Jenny Duchier et Jean-Bernard RUIDAVETS pour les nombreux bons moments passés ensemble lors des pauses et des repas.

Un remerciement également à Jennifer Cros pour ses précieuses relectures d'articles en anglais, ainsi qu'à Vincent Courjault-Radé et Florian Lefebvre pour leur soutien tout au long de cette aventure.

Je remercie de plus tous mes amis toulousains, en particulier les résidents de la Colombe, Io, Matthieu, Lorena Marie, Charlotte, Antoine et Pierrot, ainsi que Sousou et Gatsby pour tous les bons moments passés ensemble. Je tiens à remercier finalement ma famille sans qui tout cela n'aurait pas été possible pour leur soutien de longue date bien que mon univers ne leur soit pas familier, merci à vous les artistes.

Table des matières

Introduction	2
1 Principe de base de la modélisation du recrutement	10
1.1 Généralités sur les modèles poissonniens	10
1.1.1 Rappels et définitions	10
1.1.2 Modélisations des intensités de recrutement	12
1.2 Modèles Poisson-Gamma	12
1.3 Estimation	13
1.3.1 Problème du nombre de centres	13
1.3.2 Estimation du paramètre θ par maximum de vraisemblance	14
1.3.3 Distribution asymptotique de θ	15
1.4 Prédiction	17
1.4.1 Loi <i>a posteriori</i> des intensités de recrutement	17
1.4.2 Prédiction du processus d'inclusion	18
1.5 Validation de l'hypothèse Gamma - Courbes d'occupation	20
2 Enrichissement du modèle	23
2.1 Sensibilité du modèle aux paramètres de la dynamique de recrutement	23
2.1.1 Procédures de générations des données	24
2.1.2 Résultats	29
2.1.3 Recommandations	33
2.2 Dates d'ouvertures inconnues : Modèle Poisson-Gamma uniforme	37
2.2.1 Estimation par maximum de vraisemblance	37
2.2.2 Propriétés de l'estimateur	38
2.2.3 Validation du modèle	38
2.3 Prise en compte des sorties d'étude	39
2.3.1 Introduction	39
2.3.2 Modèles de prise en compte des sorties d'étude	41
2.3.3 Estimation des paramètres au temps intermédiaire t_1	42
2.3.4 Prédiction.	46
2.4 Prise en compte des interruptions de recrutement	52
2.4.1 Le modèle Poisson-gamma avec incorporation des interruptions de recrutement	53
2.4.2 Procédure de génération des données	56
2.4.3 Résultats et discussion	58
2.4.4 Conclusions	64

3	Applications du modèle	66
3.1	A l'étude de faisabilité d'un essai	66
3.1.1	Indicateurs selon les paramètres a priori d'experts	66
3.1.2	Modélisation des paramètres de la phase de recrutement	68
3.1.3	Comparaison de proximité des intensités entre centres partagés pour les deux approches.	68
3.1.4	Prediction de la durée du recrutement pour le nouvel essai - u_i connus	73
3.1.5	Résultats d'analyse	75
3.1.6	Recommandations et exemple illustratif - Faisabilité de l'essai IFM 2014	80
3.1.7	Conclusion	82
3.2	A la prédiction dynamique selon des critères de jugements sur le nombre de patients	84
3.2.1	Randomisation : impact sur l'essai	84
3.2.2	Indicateurs de performances de la dynamique de recrutement	87
3.3	Au suivi dynamique du recrutement d'un essai clinique avec outcome longitudinal	92
3.3.1	Outcome : données de survie	93
3.3.2	Essai clinique avec plusieurs visites : multi-états	94
3.4	A la modélisation du coût	95
3.4.1	Introduction	95
3.4.2	Préliminaires	97
3.4.3	Application aux essais cliniques multicentriques	99
4	Conclusion et perspectives	103
4.1	Perspectives	103
4.1.1	Perspectives appliquées	104
4.1.2	Perspectives théoriques	104

Introduction

Essai clinique

Un essai clinique (étude clinique ou encore essai thérapeutique) est une recherche biomédicale pratiquée sur l'Homme dont l'objectif est la consolidation et le perfectionnement des connaissances biologiques ou médicales. Elles portent ainsi sur des problématiques variées, notamment :

- Utilité d'un médicament pour traiter une maladie donnée.
- Comparaison de la sécurité et de l'efficacité entre une dose prédéterminée et une dose différente d'un médicament.
- Évaluation d'un nouveau dispositif de diagnostic par rapport à une méthode de référence.
- Étude de l'efficacité d'un médicament dont l'AMM (Autorisation de Mise sur le Marché) a été validée pour une autre maladie.
- Comparaison de l'efficacité entre différents traitements pour traiter une même maladie.

Ainsi lors de l'élaboration d'un nouveau médicament, plusieurs essais cliniques sont à mener selon les recherches déjà existantes : test d'efficacité et d'innocuité du traitement, étude de la dose optimale, comparaison à un traitement ayant obtenu l'AMM, comparaison à l'effet placebo.

L'élaboration de ces essais cliniques est réglementée. Ainsi pour que des données cliniques puissent permettre l'obtention d'une AMM, elles doivent suivre entre autre les Bonnes Pratiques Cliniques (BCP) décrivant une norme relative à la bioéthique dès lors que le sujet est humain. Ces dernières ont depuis 1997 été harmonisées au sein de plusieurs pays grâce à l'initiative conjointe de l'industrie pharmaceutique et des autorités réglementaires de l'Europe, du Japon et des États-Unis de par le Conseil International d'harmonisation des exigences techniques pour l'enregistrement des médicaments à usage humain créée en 1990 [30]. Les données pré-cliniques (menées en amont sur l'animal) quant à elles sont soumises aux Bonnes Pratiques de Laboratoire telles que définies par L'OCDE [51].

Le droit de la santé prévoit quant à lui l'obligation d'informer les participants des risques éventuels de façon exhaustive ainsi que l'obligation d'obtenir son consentement éclairé à l'inclusion dans l'essai. Il est également obligatoire d'avoir l'avis favorable d'un Comité de Protection des Personnes (CPP). Son but est de vérifier entre autre l'intérêt scientifique et médical de l'essai, d'évaluer le rapport entre risques et bénéfices, la conformité aux BCP ainsi que la présence d'une assurance en cas de dédommagement

de participants.

L'application d'un modèle statistique peut, selon la méthode employée, être sujette à différents biais, causant ainsi des erreurs de jugement sur la conclusion de l'essai. Les essais cliniques ne sont pas exempts de tels biais, il est donc nécessaire d'intégrer dans la méthodologie leur prise en compte.

Les biais connus lors d'essais cliniques sont celui d'attrition : différences d'effectifs entre la phase d'inclusion et la fin de l'étude pour le groupe témoin et/ou le groupe test en fonction de facteurs excluant (intolérance au traitement, effets secondaires) pouvant amener à une sur-estimation de l'efficacité ; le biais de confusion (confusion entre les effets du traitement thérapeutique et les conséquences de la maladie traitée), le biais de sélection (différence entre le groupe traité et le groupe témoin), le biais de suivi (différences de prise en charge entre le groupe traité et le groupe témoin) ainsi que le biais d'évaluation (critère de jugement déterminé de manière différente selon les deux groupes).

Pour pallier à ces sources d'erreurs potentielles, plusieurs méthodologies ont été mises en place. Un essai clinique sera autant que possible une étude contrôlée (un groupe témoin permettant de comparer l'effet du médicament présumé actif à un traitement connu inactif mais inoffensif), et randomisée (la répartition de la population entre le groupe traité et le groupe témoin doit être faite de manière aléatoire) et menée en aveugle ou en insu (voire en double aveugle dans le cas où l'expérimentateur peut être source de biais d'interprétation).

Pour le biais d'attrition il est nécessaire d'avoir une étude en intention de traiter (tout patient reste étudié dans le groupe auquel il a été assigné malgré des possibles sorties d'études) permettant de ne pas faire d'analyse per protocole.

Ainsi le développement d'un nouveau médicament en vue d'un usage thérapeutique se décompose le plus souvent en plusieurs études cliniques successives que l'on différencie en quatre phases, elles mêmes précédées d'études pré-cliniques.

Phases de développement d'un médicament

Phase pré-clinique : Lors de cette phase la molécule est tout d'abord étudiée, au point de vue de sa structure, de ses effets cellulaires, biologiques et comportementaux sur l'animal.

Elle est tout d'abord menée in vitro, puis in vivo sur animaux (premièrement sur des rongeurs, puis des non-rongeurs : des porcs (proximité génétique), chiens et primates). Ces tests permettent de déterminer les aspects pharmacocinétiques de la molécule (étude de son devenir au sein de l'organisme), ses aspects pharmacodynamiques (étude de ses effets sur l'organisme), la dose maximale tolérée chez l'animal, la dose sans effet observable et la dose sans effet toxique observable.

Si cette phase montre des effets prometteurs et inoffensifs, le médicament passe en phase I.

Phase I : Cette phase est le préliminaire pour l'étude de sécurité du médicament. Le but est de déterminer la tolérance et l'absence d'effets indésirables chez l'homme sur une population de volontaires sains de petit effectif : 12-40 patients (elle peut également

inclure des patients en impasse thérapeutique pour lesquels le traitement représente la seule chance de survie).

Ainsi est définie une dose maximale tolérée pour l'homme à partir d'une dose initiale définie à partir des résultats de la phase pré-clinique.

Elle permet également d'étudier la pharmacocinétique et le métabolisme de la substance chez l'homme.

Phase II : Cette phase est la première administration du traitement sur une population cible. Elle se subdivise en deux sous-phases :

- Phase IIa : elle sert à établir l'efficacité de la molécule sur la pathologie considérée chez l'Homme ("proof of concept"). La population considérée est de 100-200 patients.
- Phase IIb : elle permet de déterminer la dose ainsi que les conditions de prescription optimales et les possibles effets indésirables. La population considérée est cette fois de 100-400 patients.

Phase III : Également appelée étude pivot, elle permet la comparaison de l'efficacité du traitement par rapport à un traitement de référence ou un placebo. Les groupes sont de taille importante (500-3000 patients). Les patients volontaires sont sélectionnés selon des critères d'inclusion/exclusion en fonction du traitement, de la maladie, et des connaissances biomédicales établies.

C'est à l'issue de cette phase que l'AMM peut être obtenue.

Phase IV : Elle commence si l'AMM est obtenue selon les données cliniques des phases précédentes. Elle permet l'étude de possibles effets secondaires et complications tardives à long terme, la comparaison à d'autres molécules (déjà existantes ou en développement), mais également les effets réels du traitement (seul ou en association à d'autres traitements). Elle peut également permettre de déterminer des facteurs influents quant à l'efficacité ou l'apparition d'effets indésirables.

Ainsi peuvent être modifiés en conséquence le rayon d'action du médicament, l'extension d'AMM, la posologie.

Méthodologie d'un essai clinique.

Un essai clinique est ainsi décrit par un protocole strict suivant une méthodologie normalisée [50], définissant entre autre la mise en place d'une étude prospective, dont l'objectif est de spécifier les critères d'inclusion, les critères d'exclusion, la méthodologie statistique et notamment le nombre de sujets nécessaire (NSN). Le nombre de sujets nécessaire, que nous noterons N , est le nombre minimal de patients à inclure dans l'essai afin d'assurer au test statistique une puissance donnée pour pouvoir conclure. Pour des études comparatives, le calcul du NSN est obligatoire notamment pour la publication de l'étude [50].

Période de recrutement : méthodes statistiques.

La constitution de cette base de données de N patients vérifiant tous les critères d'inclusion et aucun critères d'exclusion est un des enjeux fondamentaux d'un essai clinique. Pour ce faire plusieurs centres investigateurs sont, en général, sollicités (également pour limiter le biais de selection). La période entre l'ouverture du premier centre investigateur et le recrutement du N -ième patient est appelée période de recrutement. Cette durée est bien sûr estimée avant le lancement de l'essai. Elle demeure néanmoins aléatoire et très difficile à contrôler. C'est pourtant un paramètre fondamental de l'essai pour au moins trois raisons :

- **Raison éthique.** Il est vain de continuer une étude qui se ne finira pas faute de recrutement.
- **Raison économique.** D'une part les coûts de recherche et développement d'un médicament sont très coûteux et ne cessent d'augmenter (**en moyenne de 231 millions de dollars en 1991 [26], 403 millions en 2003 [25] et 1395 millions de dollars en 2013 [24]**). D'autre part, la durée de l'essai est comprise dans la période d'exploitation de 20 ans d'un médicament [60] un délai dans le recrutement et donc un délai dans la durée de l'essai engendre une perte de revenus qui peut être substantielle. A contrario une meilleure gestion du suivi de l'essai peut engendrer des économies.
- **Raison scientifique.** Durant la durée de l'essai, une meilleure compréhension des phénomènes biologiques sous-jacents ou le développement de nouveaux médicaments peuvent être approuvés par les agences de réglementation et entraîner l'obsolescence du médicament étudié.

La question du suivi du recrutement est évidemment au centre de la problématique des essais cliniques. Les professionnels sont demandeurs d'outils rationnels de suivi et du développement d'indicateurs permettant de détecter au plus vite des comportements "hors normes" et hors recrutement (démarrages tardifs, taux d'événements indésirable, recrutement faible,...) ou bien d'indicateurs prédictifs sur les comportement de la dynamique de recrutement à venir. En général, ces indicateurs sont basés sur une hypothèse de normalité et détectent en fonction d'un seuil prédéfini des comportements inhabituels en utilisant les moyennes et écart-type de la cohorte (approche dite 1-2 ou 3 sigma) et s'expriment sous la forme :

$$X > \text{Mean}(\text{Cohort}) + K * \text{SD}(\text{Cohort}), \quad K = 1, 2, 3 \quad (1)$$

Dans le cas où les données sont en effet distribuées selon une loi gaussienne, ces indicateurs permettent d'établir des seuils au delà desquels les données seront considérées "hors-normes". Cependant dans le cas d'une dynamique de recrutement, certaines données ne suivent pas de distribution gaussienne, ainsi l'indicateur perd en pertinence. La Figure 1 représente deux types de données recueillies à partir d'une dynamique de recrutement, et souligne la nécessité de travailler avec des indicateurs différents.

La Figure 1 illustre à quel point cette hypothèse de normalité est éloignée de la réalité des données d'une dynamique de recrutement. La question d'une modélisation d'indicateurs plus pertinente se pose donc.

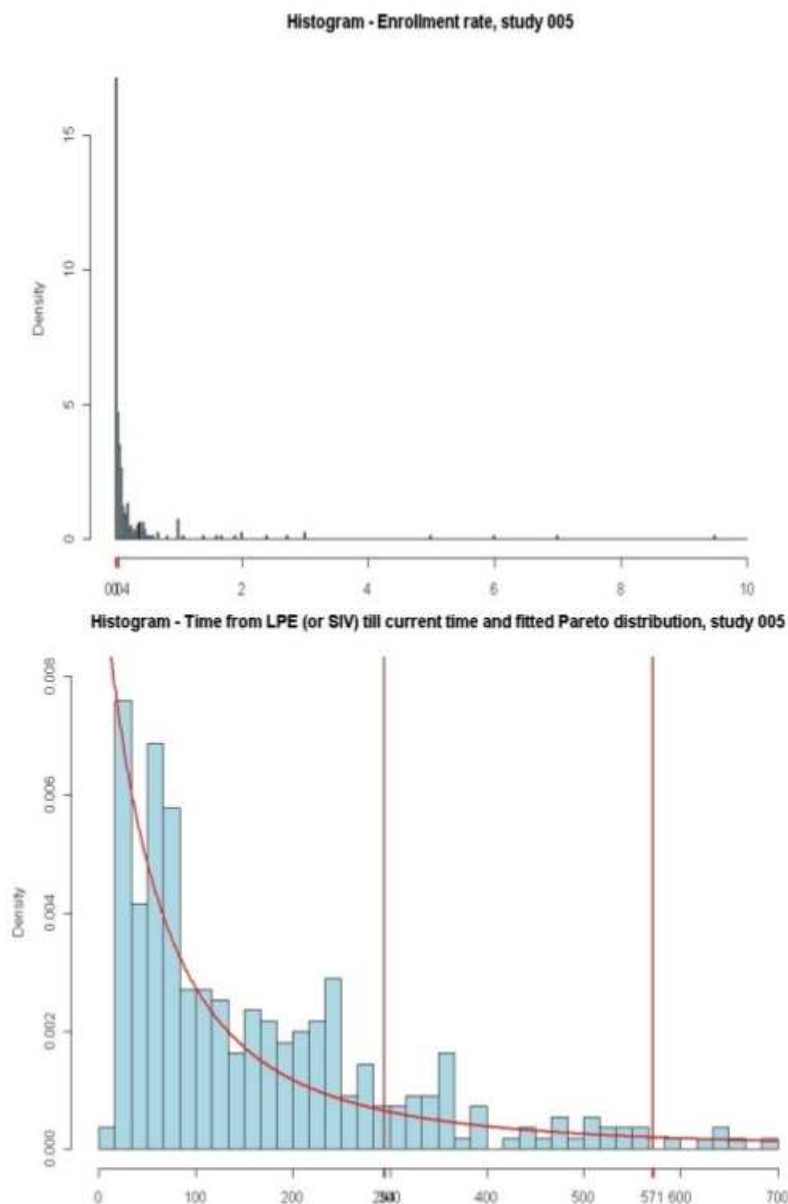


FIGURE 1 – En haut : histogramme des taux de recrutement par centres, en bas : histogramme des durées du dernier patient inclus par centre à partir de l’instant courant.

Les premières modélisations de périodes de recrutement remontent à presque 50 ans avec les travaux de Lee [33], Williford et al. [61] et Morgan [19] avec l’idée déjà d’une modélisation de la dynamique de recrutement par des processus de Poisson. Senn [57] montre que cette hypothèse est pertinente en sachant que les processus de recrutement vérifient les principales propriétés des processus de Poisson (données à valeurs entières et dont les accroissements sont indépendants et stationnaires). Un autre avantage majeur de cette hypothèse découle de l’aspect multicentrique des dynamiques étudiées. Il est en effet possible de considérer le processus de recrutement global comme la somme de l’ensemble des processus de Poisson considérés pour chaque centre. Ce processus global sera, de par la propriété d’additivité des processus de Poisson, lui aussi un processus de Poisson [56].

Un problème émerge cependant rapidement de ces modèles, à savoir la caractérisation de la moyenne et de la variance agissant sur les différentes dynamiques de recrutement

par un seul paramètre, l'intensité, introduisant un phénomène de sur-dispersion. Cette lacune est soulignée par Carter et al. [17, 18] qui préconisent l'ajout d'aléa dans les intensités de recrutement. Ces modèles sont des modèles de Poisson doublement stochastique ou processus de Cox. Des états de l'art de ces modèles statistiques ont été réalisées [15].

Le modèle qui semble être le plus pertinent à la fois d'un point de vue de la validation théorique mais aussi pratique est le modèle dit Poisson-gamma [10, 46] basé sur un processus de Poisson dont les intensités par centres sont considérées comme un échantillon de loi gamma. Ce modèle est au centre de ce mémoire, sa présentation, les enrichissements ajoutés depuis sa mise en place et de multiples applications sur cas réelles seront présentés.

Modèle Poisson-Gamma.

Le modèle considère une intensité de recrutement elle même aléatoire suivant une distribution Gamma supposée constante au cours du temps.

Tout d'abord la validité de ce modèle est établie de façon asymptotique. Il est donc naturel de se poser la question de la vitesse de convergence. Celle-ci est multifactorielle, elle dépend bien sûr du nombre de centres, du nombre de patients recrutés et de la distribution de la loi gamma. Le modèle et les validations asymptotiques sont présentées au Chapitre 1.

Certaines hypothèses ne couvrent pas l'ensemble des données émanant des dynamiques de recrutement. En effet certains phénomènes peuvent modifier les données de recrutement, comme des pauses dans le recrutement, et ainsi invalider en théorie l'application du modèle sur des données ne vérifiant plus ses hypothèses. Des pistes d'enrichissement du modèle quant à ces problématiques sont envisagées au Chapitre 2.

Une étude par simulation effectuée au paragraphe 2.1 permet de donner des informations précises sur la validité du modèle dans un cas précis.

Il sera question d'inclure dans le modèle le phénomène de sortie d'étude (patients qui sont inclus mais qui sortent de l'étude avant la randomisation) pendant la phase de screening de manière plus fine que de simplement considérer une inflation entre le nombre de patients randomisés et le nombre de patients inclus de 20 %. Une technique Bayésienne empirique analogue à celle du processus de recrutement est introduite pour modéliser les sorties d'étude. Ce modèle est décrit au paragraphe 2.3.

Dans le cas où les dates d'ouvertures des centres ne sont pas connues, l'introduction d'une variable uniforme permet de garder l'efficacité prédictive du modèle comme présenté au paragraphe 2.2.

Une autre question qui se pose est : comment prendre en compte les informations d'interruptions dans le modèle dynamique de recrutement ? Le paragraphe 2.4 étudie cette question ainsi que l'apport ou non de la prise en compte de ces interruptions qui n'améliore pas forcément les performances prédictives du modèle. Il sera aussi question de mesurer l'effet de ces interruptions de recrutement sur les estimations du modèle prédictif.

Ce modèle est utilisé depuis une dizaine d'années donnant des prédictions fiables. De multiples applications sont désormais disponibles comme nous le verrons au Chapitre

3.

Il est en effet possible d'effectuer des études de faisabilité présentées au paragraphe 3.1. La dynamique de recrutement a une influence sur les caractéristiques intrinsèques au recrutement (nombre de centres vides, probabilités de recrutement, nombre de patients en cours d'étude, à chaque visite). Ces aspects fondamentaux fournissent des indicateurs pertinents pour le suivi des essais. Ils seront présentés au paragraphe 3.2.

La problématique de la dynamique de recrutement peut être couplée avec la dynamique de l'étude en elle-même quand celle-ci est longitudinale. L'indépendance des deux processus permet une estimation facile des différents paramètres. Le résultat est un modèle global du parcours du patient dans l'essai. Deux exemples clés de telles situations sont les données de survie - la modélisation permet alors d'estimer la durée d'un essai quand le critère d'arrêt est le nombre d'événements observés et les modèles de Markov - la modélisation permet alors d'estimer le nombre de patients dans un certain état au bout d'un certain temps. Ces deux exemples seront approfondis au paragraphe 3.3.

Il peut également être considéré d'emboîter ce modèle avec un modèle de suivi de l'état des patients lorsque la phase de recrutement et la phase de suivi se chevauchent (cf paragraphe 3.3.2). Enfin un modèle de coût de l'essai clinique a été introduit au paragraphe 3.4 et étudié dans [45].

Chapitre 1

Principe de base de la modélisation du recrutement

La modélisation présentée dans ce chapitre est une modélisation par approche bayésienne de l'enrôlement de patients dans un essai multi-centriques, modèle développé par Anisimov [10]. L'idée est de considérer un processus de Cox dont les intensités sont distribuées selon une loi Gamma. Le modèle est appelé modèle Poisson-gamma.

L'objet de ce chapitre sera, après un rappel sur les processus de Poisson, de présenter ce modèle de manière théorique et d'introduire les estimations des paramètres du modèle à partir de données recueillies à un instant intermédiaire de l'essai.

1.1 Généralités sur les modèles poissonniens

La modélisation de la dynamique de recrutement repose actuellement sur un modèle poissonnien du recrutement de chacun des centres. Cette idée est en effet naturelle en analogie avec la théorie des files d'attente dont le recrutement de patients peut être vu comme un exemple. Les processus de Poisson apparaissent ainsi comme naturels dans la littérature [56].

En effet en se référant aux notations de Kendall dans le cadre de processus de file d'attente, on peut associer certains de ces paramètres à des paramètres de notre modèle.

Ainsi la loi de distribution des temps entre les arrivées des clients Λ devient la loi des temps entre deux recrutements, le tampon ("buffer") devient notre population cible, le processus de sortie d'étude est associé au phénomène de drop-out, et le processus d'entrée devient le processus de recrutement.

Il est donc logique de modéliser le temps de recrutement par un processus de Poisson comme en théorie de file d'attente.

1.1.1 Rappels et définitions

Considérons un essai clinique impliquant C centres investigateurs et pour lequel le recrutement de N patients est nécessaire pour avoir la puissance de test voulue.

Definition Soit λ une fonction intégrable. Un processus stochastique $\mathcal{N}$ est un processus de Poisson d'intensité λ si ses accroissements sont indépendants et s'il possède

une distribution de la forme :

$$\forall q < p, \quad \mathcal{N}_p - \mathcal{N}_q \stackrel{\mathcal{L}}{=} \mathcal{P} \left(\int_q^p \lambda(s) ds \right)$$

Proposition 1 Soit $\mathcal{N}$ un processus de Poisson d'intensité λ partant de 0. Alors à un temps t , il suit une loi de Poisson de paramètre λt , et sa moyenne et sa variance sont données en t par :

$$\mathbb{E}[\mathcal{N}_t] = \lambda t, \quad \mathbb{V}[\mathcal{N}_t] = \lambda t$$

Cette définition peut être étendue (comme effectué lors de l'utilisation du modèle Poisson-gamma) à des intensités λ qui sont considérées aléatoires, on parle alors de processus de Poisson conditionnel ou processus de Cox.

Définition Soit $\mathcal{N}$ un processus de Poisson d'intensité λ . On appelle $\mathcal{N}$ un processus de Cox si $\{\lambda(t), t \geq 0\}$ est lui même un processus stochastique.

Proposition 2 Soit $\mathcal{N}$ un processus de Cox d'intensité λ . Alors pour tout $t \geq 0$, $\mathcal{N}_t$ a pour moyenne et variance :

$$\mathbb{E}[\mathcal{N}_t] = \mathbb{E}[\mathbb{E}[\mathcal{N}_t | \lambda]] = \mathbb{E}[\lambda]t, \quad \mathbb{V}[\mathcal{N}_t] = \mathbb{V}[\mathbb{E}[\mathcal{N}_t | \lambda]] + \mathbb{E}[\mathbb{V}[\mathcal{N}_t | \lambda]] = \mathbb{V}[\lambda]t^2 + \mathbb{E}[\lambda]t.$$

Par la suite les dynamiques de recrutement de chacun des C centres seront modélisées par des processus de Cox.

Notons $(\mathcal{N}^i, i = 1, \dots, C)$ l'ensemble des processus de Cox modélisant l'ensemble des dynamiques de recrutement de l'essai, et $\mathcal{L}(i, \theta)$ les distributions des intensités de ces processus $\mathcal{N}^i$.

Ces intensités peuvent être dépendantes du temps t (selon des fonctions plus ou moins complexes), le cas le plus fréquent lors d'étude réelle est la disparité des dates d'ouvertures entre les C centres.

Notons $(u_i, i = 1, \dots, C)$ l'ensemble de ces dates d'ouvertures, les intensités de recrutement (toujours considérées constantes à partir du temps u_i) sont alors égales à :

$$\lambda_i(t) = \lambda_i \chi(u_i \leq t),$$

où $\chi(A) = 1$ si A est vrai et $\chi(A) = 0$ sinon.

Proposition 3 La somme de deux processus de Cox d'intensités respectives λ et ψ est un processus de Cox d'intensité $\lambda + \psi$

Notons $\mathcal{N}^C$ le processus global défini comme $\mathcal{N}^C = \sum_{i=1}^C \mathcal{N}^i$. C'est un processus de Cox par la proposition 3 d'intensité $\Lambda = \sum_{i=1}^C \lambda_i$.

La propriété d'additivité est très importante dans l'étude des dynamiques de recrutement multicentriques car elle permet de considérer la dynamique globale (somme des C dynamiques de recrutement particulières) dont l'intensité est la somme des C intensités spécifiques à chaque processus $\mathcal{N}^i$.

La variable d'intérêt du modèle, à savoir la variable modélisant le temps d'atteinte des N patients se définit comme $\mathcal{T}^C = \inf_{t \geq 0} \{\mathcal{N}^C(t) \geq N\}$.

1.1.2 Modélisations des intensités de recrutement

Le choix des hypothèses posées sur les intensités de recrutement dans le cas d'une dynamique modélisée par des processus de Poisson est une question fondamentale pour la pertinence de ce modèle.

La flexibilité du modèle est en effet en grande partie assurée par les valeurs et estimations des λ_i . Plusieurs situations existantes sont présentées :

- On fait confiance aux λ_i fournis par les investigateurs. Dans ce cas le paramètre est une constante et les résultats découlent de calculs élémentaires sur la loi de Poisson.
- On s'inspire des paramètres issus de données historiques. Ce problème tout à fait fondamental est décrit au paragraphe 3.1.
- On s'inspire des résultats sur le recrutement obtenus à un instant intermédiaire T_1 . La encore plusieurs options peuvent être envisagées :
 - On estime chacun des λ_i par le rapport entre le nombre de patients recrutés et la durée d'activité du centre. Cette approche est à proscrire pour deux raisons, tout d'abord, quand le nombre de centre C est grand, le nombre d'estimations est grand et leur pertinence douteuse. Ensuite cette méthode manque de variabilité, en effet si l'on considère un centre qui n'a pas encore recruté de patients, l'estimation de son intensité est 0 et ce centre n'est donc pas autorisé à recruter plus tard.
 - On considère les λ_i comme un échantillon de taille C distribué selon une certaine loi $\mathcal{L}(\theta)$ seuls les paramètres θ sont à estimer ce qui résout les deux problèmes précédents.

Cette dernière modélisation, se basant sur une approche bayésienne empirique, est en pratique le plus convaincant [10] notamment quand la distribution *a priori* est gamma [46]. Par la suite nous nous placerons dans ce cadre.

1.2 Modèles Poisson-Gamma

Cette méthode est qualifiée d'approche bayésienne empirique du fait de l'estimation de l'hyperparamètre θ de la famille paramétrique considérée par les observations et non par une attribution d'une nouvelle loi *a priori*.

En effet il est supposé que les intensités de recrutement sont distribuées selon une certaine loi *a priori* appartenant à une famille paramétrique $\{p_\theta, \theta \in \Theta\}$. Ce paramètre θ est ensuite estimé par $\hat{\theta}$ par maximum de vraisemblance. Les lois *a posteriori* des intensités de recrutement seront ensuite calculées en utilisant cette estimation $\hat{\theta}$, contrairement à une méthode bayésienne stricte où une nouvelle loi *a priori* est attribué au paramètre θ .

D'autre famille paramétrique ont été étudiées comme la loi Pareto [44] mais la simplicité de l'utilisation d'une famille paramétrique gamma sans perte d'efficacité du modèle prédictif motive son utilisation.

Par la suite l'écriture $a \vee b$ désignera $\max(a, b)$. Voici quelques définitions nécessaires à la clarification du modèle Poisson-gamma.

- C désigne le nombre de centres.
- N le nombre de sujets cible de l'étude.
- u_i désigne la date d'ouverture du i ème centre (si elles sont disponibles).
- t_1 désigne le temps intermédiaire d'étude.
- k_i désigne le nombre de patients recrutés par le centre i au temps t_1 .
- τ_i désigne la durée d'activité du centre i jusqu'à t_1 ($\tau_i = (t_1 - u_i) \vee 0$).
- λ_i est l'intensité de recrutement du centre i modélisé par un processus de Poisson $\{N_i(t), t \geq t_1\}$.
- $(\lambda_i, i = 1, \dots, C)$ est considéré comme un échantillon indépendant de variables aléatoires de loi gamma de paramètres (α, β) .

Il est à noter que si X est une variable aléatoire de loi $\Gamma(\alpha, \beta)$ dont la densité est donnée par :

$$p_{\alpha, \beta}(x) = \kappa e^{-\beta x} x^{\alpha-1} \chi(x > 0),$$

alors,

$$\mathbb{E}[X] = \frac{\alpha}{\beta}; \quad \mathbb{V}[X] = \frac{\alpha}{\beta^2}.$$

Il est maintenant possible de modéliser les intensités *a posteriori*. Les paramètres du modèle seront recalibrés selon les données collectées par inférence bayésienne.

1.3 Estimation

1.3.1 Problème du nombre de centres

Le nombre de centres impliqués dans la dynamique de recrutement est un paramètre fondamental pour la décision du modèle de prédiction à utiliser, les différents cas de figures sont présentés ci-après.

Ce nombre est forcément dépendant de plusieurs autres paramètres, une étude par simulation a été effectuée pour étudier ce phénomène appliqué au modèle Poisson-gamma au paragraphe 2.1. Une première approche introductive est présentée ci dessous.

Dynamique multicentrique : peu de centres

Lorsque C est 'petit', on peut modéliser la dynamique de recrutement de chaque centre sans forcément supposer des intensités aléatoires et appliquer la proposition 3 pour obtenir le modèle complet. Ceci est pertinent car le phénomène de sur dispersion est minimisé du fait du petit nombre de centres.

Dynamique multicentrique : beaucoup de centres

Lorsque le nombre de centres C augmente pour un nombre N fixe de patients requis, la non prise en compte de la variabilité par un processus de Cox implique des biais d'estimations forts qui peuvent grandement impacter la pertinence de la prédiction.

De plus, que ce soit dans le cas de processus de Poisson ou de Cox dont chaque λ_i peut être considéré certes aléatoire mais pouvant être potentiellement de loi différentes, le nombre d'estimations nécessaires augmente grandement.

C'est pourquoi l'utilisation de Processus de Cox dans le cas de dynamique de recrutement multicentrique avec C assez grand est largement utilisé en considérant entre autre que l'ensemble $(\lambda_i, i = 1, \dots, C)$ est une réalisation d'une variable aléatoire dont les paramètres de la loi sont à estimer. Ces hypothèses permettent premièrement une meilleure prise en compte de la variabilité du problème en réduisant considérablement le nombre d'estimations nécessaires. Le modèle Poisson-Gamma fait partie de cette catégorie de modélisation en supposant le vecteur des λ_i comme une réalisation d'une variable aléatoire suivant une loi $\Gamma(\alpha, \beta)$.

1.3.2 Estimation du paramètre θ par maximum de vraisemblance

Les paramètres a priori de la loi Gamma utilisée dans le modèle sont rarement connus ou estimés à l'avance. Les techniques de calibrage du modèle par estimation empirique bayésienne sont donc utilisées grâce aux données collectées jusqu'à t_1 . La méthode d'estimation par maximum de vraisemblance sera la méthode d'estimation utilisée par la suite.

Remarque Anisimov et al. [10] ont utilisé d'autres méthodes d'estimations du paramètre θ comme la méthode des moindres carrés ou encore la méthode des moments. Cependant devant la similarité des résultats, la méthode d'estimation par maximum de vraisemblance est utilisée sachant qu'elle nous permet de calculer également la variance asymptotique de l'estimateur.

La loi du vecteur $(k_1, \dots, k_C)$ servant au calcul de la vraisemblance est égale par indépendance des processus d'inclusion $(\mathcal{N}_i, i = 1, \dots, C)$ à [10] :

$$\mathbb{P}[k_1 = n_1, \dots, k_C = n_C] = \prod_{i=1}^C \frac{\Gamma(\alpha + n_i)}{n_i! \Gamma(\alpha)} \frac{\beta^\alpha \tau_i^{n_i}}{(\beta + \tau_i)^{\alpha + n_i}}.$$

Sachant cela, nous pouvons maintenant estimer le paramètre $\theta = (\alpha, \mu)$

Théorème 1 ([7]) *L'estimateur du maximum de vraisemblance $\hat{\theta} = (\hat{\alpha}, \hat{\mu})$ des paramètres (α, μ) est obtenu par maximisation de la log-fonction :*

$$M_C^\Gamma(\alpha, \mu) = \alpha \ln(\alpha/\mu) - \ln \Gamma(\alpha) + \frac{1}{C} \sum_{i=1}^C [\ln \Gamma(\alpha + k_i) - (\alpha + k_i) \ln(\alpha/\mu + \tau_i)].$$

Preuve Lorsque l'intensité λ d'un processus de Poisson est non déterministe, la fonction de densité d'un processus de Poisson f est remplacée par $\mathbb{E}_\lambda[f]$ [46].

Sachant cela :

$$\begin{aligned}
 \mathbb{P}[\mathcal{N}_i(t_1) = n_i] &= \mathbb{E}_{p_{\alpha,\beta}} [\mathbb{P}[\mathcal{N}_i(t_1) = n_i]] \\
 &= \mathbb{E}_{p_{\alpha,\beta}} \left[\exp\left(-\int_0^{t_1} \lambda_i(s) s ds\right) \frac{\left(\int_0^{t_1} \lambda_i(s) s ds\right)^{n_i}}{n_i!} \right] \\
 &= \mathbb{E}_{p_{\alpha,\beta}} \left[\exp(-\lambda_i(t_1 - u_i)) \frac{(\lambda_i(t_1 - u_i))^{n_i}}{n_i!} \right] \\
 &= \mathbb{E}_{p_{\alpha,\beta}} \left[\exp(-\lambda_i \tau_i) \frac{(\lambda_i \tau_i)^{n_i}}{n_i!} \right] \\
 &= \frac{\tau_i^{n_i}}{n_i!} \int_0^\infty \exp(-x \tau_i) x^{n_i} p_{\alpha,\beta}(x) dx
 \end{aligned}$$

Par indépendance entre les processus $\mathcal{N}_i$ pour $i = 1 : C$, on obtient :

$$\begin{aligned}
 \mathbb{P}[\mathcal{N}_1(t_1) = n_1, \dots, \mathcal{N}_C(t_1) = n_C] &= \prod_{i=1}^C \frac{\tau_i^{n_i}}{n_i!} \int_0^\infty \exp(-x \tau_i) x^{n_i} p_{\alpha,\beta}(x) dx \\
 &= \prod_{i=1}^C \frac{\Gamma(\alpha + n_i)}{n_i! \Gamma(\alpha)} \frac{\beta^\alpha \tau_i^{n_i}}{(\beta + \tau_i)^{\alpha + n_i}}
 \end{aligned}$$

Par conséquent,

$$\begin{aligned}
 M_C^F(\alpha, \mu) &= \log(\mathbb{P}[\mathcal{N}_1(t_1) = n_1, \dots, \mathcal{N}_C(t_1) = n_C]) \\
 &= \alpha \ln(\alpha/\mu) - \ln \Gamma(\alpha) + \frac{1}{C} \sum_{i=1}^C [\ln \Gamma(\alpha + k_i) - (\alpha + k_i) \ln(\alpha/\mu + \tau_i)].
 \end{aligned}$$

Trois remarques poussent à préférer l'écriture de ces résultats en fonction de (α, μ) au lieu de (α, β) , où $\mu = \alpha/\beta$:

- D'un point de vue pratique, le paramètre μ désigne directement la moyenne de la loi gamma considérée.
- L'estimation du paramètre μ est plus stable que celle de β comme souligné dans [6].
- Cette reparamétrisation permet de travailler avec une matrice d'information globale de Fisher qui est diagonale, réduisant le nombre de termes à calculer.

1.3.3 Distribution asymptotique de θ

Le calcul de la matrice de Fisher à partir de la fonction vraisemblance comme effectué dans [46] permet l'étude asymptotique de ces estimateurs.

En notant,

$$f(\theta, k_i, \tau_i) = \alpha \ln(\alpha/\mu) - \ln \Gamma(\alpha) + [\ln \Gamma(\alpha + k_i) - (\alpha + k_i) \ln(\alpha/\mu + \tau_i)].$$

La fonction de vraisemblance s'écrit à une constant additive près comme :

$$M_C^\Gamma = \frac{1}{C} \sum_{i=1}^C f(\theta, k_i, \tau_i),$$

La matrice d'information de Fisher d'un centre i , $J_i(\alpha, \mu)$, est de la forme :

$$\begin{pmatrix} (J_i)_{11} & 0 \\ 0 & (J_i)_{22} \end{pmatrix},$$

où

$$\begin{aligned} (J_i)_{11} &= \mathbb{E}_{(\alpha, \mu)}[\psi^{(1)}(\alpha) - \psi^{(1)}(\alpha + k_i)] - \frac{\mu \tau_i}{\alpha(\alpha + \mu \tau_i)}, \\ (J_i)_{22} &= \frac{\alpha \tau_i}{\mu(\alpha + \mu \tau_i)}, \end{aligned}$$

avec $\psi^{(1)}$ la fonction trigamma définie pour tout $x > 0$ par :

$$\psi^{(1)}(x) = \frac{d^2}{dx^2} \ln(\Gamma(x)) \quad \text{with} \quad \Gamma(x) = \int_0^\infty t^{x-1} e^{-t} dt.$$

La matrice d'information globale est, quant à elle, égale à :

$$F^C = \sum_{i=1}^C J_i(\alpha, \mu),$$

et la matrice asymptotique définie comme (dans le cas où la limite existe) :

$$F = \lim_{C \rightarrow \infty} \frac{1}{C} F^C.$$

Plusieurs hypothèses permettent de s'assurer de l'existence de la limite introduite précédemment.

Ainsi une condition suffisante pour que $\lim_{C \rightarrow \infty} \frac{1}{C} F^C$ existe et soit définie positive est que la suite $(\tau_i)_{i=1, \dots, C}$ soit uniformément distribuée dans un intervalle $[\tau_m, \tau_M]$ avec $\tau_m \geq 0$ et $\tau_M > 0$. Sachant que pour toute famille de fonctions Riemman-intégrables $(g_i, i = 1, \dots, C)$ $\lim_{C \rightarrow \infty} \frac{1}{C} \sum_{i=1}^C g_i(\tau_i) = \lim_{C \rightarrow \infty} \int_{\tau_m}^{\tau_M} \sum_{i=1}^C g_i(\tau) d\tau$. Alors :

$$F = \lim_{C \rightarrow \infty} F^C = \lim_{C \rightarrow \infty} \int_{\tau_m}^{\tau_M} \sum_{i=1}^C J_i(\tau) d\tau$$

Sachant qu'il existe $i \in \{1, \dots, C\}$, tel que $\tau_i > 0$ (sinon l'application du modèle est inutile : aucun centre n'est encore ouvert et les seules informations utilisées ici viennent de l'essai uniquement et non de données historiques externes comme présenté en Section 3.1), que par définition pour tout $i \in \{1, \dots, C\}$, $\tau_i \geq 0$ et que les fonctions $J_i(\tau)$ sont continues et définies positives, alors F existe et est définie positive.

Ainsi la fonction de covariance de la distribution asymptotique du paramètre $\theta = (\alpha, \mu)$ est $[F]^{-1}$.

La région de confiance de niveau asymptotique $(1 - \delta)$ de (α, μ) est l'ensemble :

$$\left\{ \theta \in \mathbb{R}^2; (\theta - (\alpha, \mu))^T F (\theta - (\alpha, \mu)) \leq \frac{\chi_2^2(\delta)}{C} \right\}, \quad (1.1)$$

où $\chi_2^2(\delta)$ est le δ -quantile d'une distribution du χ^2 à deux degré de liberté.

Ainsi selon la théorie de la méthode par maximum de vraisemblance, pour C assez grand et sous des conditions standard de régularité, la distribution de l'estimateur $\hat{\theta}$ est approximé par une distribution gaussienne bivariée de moyenne $(\alpha_{reel}, \mu_{reel})$ et de matrice de variance covariance F [10], nous notons cette variable B .

Ce résultat s'écrit sous la forme [10] :

$$\hat{\theta} \approx \theta + \frac{B}{\sqrt{C}} \quad .$$

1.4 Prédiction

1.4.1 Loi *a posteriori* des intensités de recrutement

On note $(\tilde{\mathcal{N}}_i, i = 1, \dots, C)$ l'ensemble des processus de Cox modélisant pour chaque centre i l'inclusion des patients après l'instant intermédiaire t_1 par ce centre, et $\tilde{\mathcal{N}}$ le processus global. Ils se définissent $\forall t \geq 0$ comme suit :

$$\tilde{\mathcal{N}}_i(t) = \mathcal{N}_i(t + t_1) - \mathcal{N}_i(t_1), \quad \text{et } \tilde{\mathcal{N}}(t) = \sum_{i=1}^C \tilde{\mathcal{N}}_i(t),$$

Soit également l'intensité globale de la modélisation de la dynamique de recrutement $\Lambda = \sum_{i=1}^C \lambda_i$. Il s'agit d'un décalage de longueur t_1 des processus pré-définis $\mathcal{N}_i$ (le recrutement avant t_1 est déjà effectué, on ne modélise que la prédiction après t_1). Le temps d'atteinte à estimer sera ainsi noté $\tilde{T} = \inf_{t \geq 0} \{ \tilde{\mathcal{N}}(t) = \tilde{n} \}$.

Comme vu précédemment, le vecteur des intensités de recrutement $(\lambda_i, i = 1, \dots, C)$ (des processus $\mathcal{N}_i$ et $\tilde{\mathcal{N}}_i$) est considéré, *a priori*, distribué selon une loi $\Gamma(\alpha, \alpha/\mu)$ (où aucune hypothèse n'est formulée sur ces paramètres mais qui seront substitués lors de la phase de prédiction du modèle par $\hat{\theta} = (\hat{\alpha}, \hat{\mu})$ calculé grâce au théorème 1).

On suppose que les dates d'ouvertures des centres $(u_i)_{1 \leq i \leq C}$ sont connues, et que les intensités d'inclusion sont distribuées suivant une loi de densité p_θ , où $\theta \in \mathbb{R}^d$.

Proposition 4 *La densité de λ_i sachant $N_i(t_1) = k_i$, notée $p_\theta^{t_1}(\cdot)$ peut être évaluée par la formule de Bayes*

$$\begin{aligned} p_\theta^{t_1}(x) &= \frac{\mathbb{P}[N_i(t_1) = k_i \mid \lambda_i = x] p_\theta(x)}{\mathbb{P}[N_i(t_1) = k_i]} \\ &= e^{-\tau_i x} x^{k_i} p_\theta(x) \frac{\tau_i^{k_i}}{k_i! \mathbb{P}[N_i(t_1) = k_i]}, \end{aligned} \quad (1.2)$$

où p_θ est la densité de λ_i et $\tau_i = t_1 - u_i$.

Dans le cas du modèle Poisson-gamma, $p_\theta(x) = \beta^\alpha \Gamma(\alpha)^{-1} e^{-\beta x} x^{\alpha-1} \chi(x \geq 0)$, et on déduit de (1.2) la proposition suivante tirée de [10].

Proposition 5 (Loi prédictive) *Dans le modèle Gamma-Poisson, la densité de λ_i sachant $N_i(t_1) = k_i$ est celle d'une loi Gamma de paramètres $(\alpha + k_i, \beta + \tau_i)$.*

Preuve La proposition 4 appliquée à la densité $p_{(\alpha, \beta)} : x \mapsto \beta^\alpha \Gamma(\alpha)^{-1} e^{-\beta x} x^{\alpha-1} \chi(x > 0)$ implique que la densité $p_{(\alpha, \beta)}^{t_1}$ de λ_i sachant k_i est de la forme

$$p_{(\alpha, \beta)}^{t_1}(x) = M e^{-\tau_i x} x^{k_i} e^{-\beta x} x^{\alpha-1} \chi(x > 0) = M e^{-(\beta + \tau_i)x} x^{k_i + \alpha - 1} \chi(x > 0),$$

où M est une quantité ne dépendant pas de x . On en déduit le résultat.

Cependant comme énoncé au paragraphe 1.1.1, ces intensités peuvent être dépendantes du temps.

Ainsi dans le cas théorique où pour tout $i \in \{1, \dots, C\}$, $u_i = 0$ alors on a bien pour chaque processus $\tilde{\mathcal{N}}_i$ une intensité de recrutement qui est égale à λ_i lui même distribué selon une loi $\Gamma(\hat{\alpha}, \hat{\alpha}/\hat{\mu})$ à partir de t_1 .

Dans le cadre du modèle où les u_i ne sont pas nécessairement nuls, les intensités de recrutement des processus d'inclusion sont égales à :

$$\lambda_i(t) = \lambda_i \chi(t_i^{t_1} \leq t),$$

où $t_i^{t_1} = t_1 \vee u_i$ et λ_i est distribué selon une loi $\Gamma(\alpha + k_i, \alpha/\mu + \tau_i)$.

Il est maintenant possible d'étudier les indicateurs de performance des centres, ainsi que d'effectuer la prédiction souhaitée.

1.4.2 Prédiction du processus d'inclusion

La phase de prédiction du modèle Poisson-gamma consiste en l'utilisation des paramètres calculés précédemment pour la prédiction de la variable d'intérêt, à savoir ici la prédiction de la durée réelle de la dynamique globale de recrutement. Par la suite nous notons $n_1 = \mathcal{N}(t_1)$ le nombre total de patients recrutés au temps t_1 .

On rappelle qu'une distribution de Pearson IV notée $P_{\text{VI}}(n, a, b)$ possède une fonction de densité :

$$p_{n,a,b}(x) = \frac{1}{\mathcal{B}(n, a)} \frac{x^{n-1} b^a}{(x + b)^{n+a}},$$

où $\mathcal{B}(a, b) = \int_0^1 x^{a-1} (1-x)^{b-1} dx$ est la fonction Beta.

Proposition 6 [36] *Soit $\tilde{\mathcal{N}}$ un processus de Cox dont l'intensité homogène Λ est distribuée selon une loi $\Gamma(\mathbf{A}, \mathbf{B})$, alors la loi de $\tilde{T}$ admet comme densité par rapport à la mesure de Lebesgue la fonction :*

$$p_{\tilde{T}} : x \rightarrow \frac{\mathbf{B}^{\mathbf{A}} \Gamma(\tilde{n} + \mathbf{A})}{\Gamma(\mathbf{A}) \Gamma(\tilde{n})} \frac{x^{\tilde{n}-1}}{(x + \mathbf{B})^{\tilde{n} + \mathbf{A}}}.$$

Sachant que $\mathcal{B}(a, b) = (\Gamma(a)\Gamma(b))/\Gamma(a+b)$, on reconnaît dans la fonction de densité précédente une loi de Pearson IV.

Pour la prédiction plusieurs cas sont à considérer, dans le cas particulier où tous les centres commencent à recruter en même temps à $t = 0$ la prédiction est simple. Si les durée d'activités sont différentes, on a le théorème suivant :

Théorème 2 ([7, 6]) *En notant $\tilde{n} = n - n_1$ le nombre de patients restant à recruter après t_1 et $\tilde{T} = \inf_{t \geq 0} \{\tilde{N}(t) = \tilde{n}\}$ la durée de recrutement restant.*

- *En supposant que pour tout $i \in \{1, \dots, C\}, u_i = 0$, (donc $\tau_i = t_1$ pour tout $i \in \{1, \dots, C\}$).*

Alors $\tilde{T}$ est distribué selon une loi $\Gamma(\tilde{n}, \Lambda)$ (donc doublement stochastique car Λ est une variable aléatoire).

Ainsi par la proposition 6 nous avons $\tilde{T}$ qui est distribué selon une loi $P_{VI}(\tilde{n}, \hat{\alpha}C + n_{t_1}, \alpha/\mu + t_1)$.

- *En supposant que pour tout i dans $\{1, \dots, C\} u_i = u > 0$, (donc $\tau_i = \tau = t_1 - u$, pour tout i).*

on a $\tilde{T}$ qui est distribué selon une loi $P_{VI}(\tilde{n}, \hat{\alpha}C + n_1, \hat{\beta} + \tau)$.

- *Dans le cas concret, les τ_i peuvent différer les uns des autres.*

La distribution de $\tilde{T}$ peut être approximée pour C assez grand ($C \geq 10$ [6]) par une distribution $P_{VI}(\tilde{n}, A, B)$ avec :

$$\mathbf{A} = \frac{\left(\sum_{i=1}^C \hat{m}_i\right)^2}{\sum_{i=1}^C \hat{b}_i^2}, \quad \mathbf{B} = \frac{\sum_{i=1}^C \hat{m}_i}{\sum_{i=1}^C \hat{b}_i^2} \quad \text{où} \quad \hat{m}_i = \frac{\alpha + k_i}{\beta + \tau_i} \quad \text{et} \quad \hat{b}_i^2 = \frac{\alpha + k_i}{(\beta + \tau_i)^2}.$$

Preuve Sachant que l'intensité globale du processus $\tilde{\mathcal{N}}$ est $\Lambda = \sum_{i=1}^C \lambda_i \sim \sum_{i=1}^C \Gamma(\hat{\alpha} + k_i, \hat{\beta} + \tau_i)$, on peut représenter le temps restant comme $\tilde{T} \sim \Gamma(\tilde{n}, 1)/\Lambda$.

Dans le cas où pour tout $i \in \{1, \dots, C\}, u_i = u$, on a par additivité de lois Gamma ayant la même variance :

$$\begin{aligned} \tilde{T} &\sim \Gamma(\tilde{n}, 1) / \Gamma\left(\sum_{i=1}^C \hat{\alpha} + k_i, \hat{\beta} + \tau\right) \\ &\sim \Gamma(\tilde{n}, 1) / \Gamma(\text{Calpha} + n_{t_1}, \text{beta} + \tau) \end{aligned}$$

Sachant que une variable $\Gamma(a, b) \sim \Gamma(a, 1)/b$, on a [36, p. 381] :

$$\tilde{T} \sim P_{VI}\left(\tilde{n}, \hat{\alpha}C + n_1, \hat{\beta} + \tau\right)$$

Dans le cas où au moins un des u_i diffère des autres, alors pour $C \geq 10$, Λ peut être approché par une variable gamma de moyenne A/B et de variance A/B^2 [7, 6] comme défini précédemment.

Ainsi $\Lambda \sim \Gamma(A, B)$, et en utilisant la proposition 6, on trouve :

$$\tilde{T} \sim P_{VI}(\tilde{n}, A, B)$$

En considérant une même écriture du triplet de paramètres de la loi de Pearson IV par $(\tilde{n}, \mathbf{A}, \mathbf{B})$ pour les trois cas du Théorème 2 envisagés, on a le théorème pour $T = t_1 + \tilde{T}$:

Théorème 3 [44]

L'estimation de la durée de recrutement résiduelle est

$$\mathbb{E}[\tilde{T}] = \begin{cases} t_1 + \tilde{n} \frac{B_1}{A_1 - 1} & \text{si } A_1 > 1 \\ +\infty & \text{si } 0 \leq A_1 \leq 1. \end{cases}$$

- La probabilité de finir à l'instant T_R est donnée par :

$$P[\tilde{T} \leq T_R - t_1] = \frac{\Gamma(A + \tilde{n})}{\Gamma(A)(\tilde{n} - 1)!} \beta_{inc} \left(\frac{T_R - t_1}{T_R - t_1 + B}; \tilde{n}, A \right).$$

où β_{inc} désigne la fonction beta incomplète définie par :

$$\beta_{inc}(x; a, b) = \int_0^x t^{a-1} (1-t)^{b-1} dt$$

Preuve Sachant Λ , la densité de $\tilde{T}$ est $t \rightarrow \frac{\Lambda^{\tilde{n}}}{(\tilde{n}-1)!} e^{-\Lambda t} t^{\tilde{n}-1}$. En conditionnant, on en déduit la densité de $\tilde{T}$:

$$\begin{aligned} p_{\tilde{T}(t)} &= \frac{B^A}{\Gamma(A)(\tilde{n} - 1)!} t^{\tilde{n}-1} \int_0^{+\infty} e^{-Bx} x^{A-1} x^{\tilde{n}e^{-xt}} dx \\ &= \frac{B^A \Gamma(A + \tilde{n})}{\Gamma(A)(\tilde{n} - 1)!} \frac{t^{\tilde{n}-1}}{(t + B)^{\tilde{n}+A}}. \end{aligned}$$

Les relations en découlent.

1.5 Validation de l'hypothèse Gamma - Courbes d'occupation

Dans une optique d'application et de validation du modèle, une étude des distributions des propriétés de l'occupation des centres a été réalisée [10]. Il sera question d'étudier la distribution du nombre de centre en fonction du nombre de patients recrutés.

On considère que tous les λ_i sont donnés. On note :

$$p_i = \frac{\lambda_i}{\Lambda}, \quad i = 1, \dots, C.$$

Ainsi le vecteur $(n_1, \dots, n_C)$ a une distribution multinomiale de paramètres $(p_1, \dots, p_C)$:

$$\mathbb{P}[n_1 = k_1, \dots, n_C = k_C] = \frac{n!}{k_1! \dots k_C!} \prod_{i=1}^C p_i^{k_i},$$

où $n = \sum_{i=1}^C k_i$.

Soient C et n donnés, nous notons $v(n, C, j)$ le nombre de centres ayant recruté exactement j patients, $j = 1, \dots, n$. Les variables $v(n, C, j)$ sont dépendantes et $C = \sum_{j=1}^n v(n, C, j)$.

Dans le cas où $p_i = 1/C$, l'étude revient à une étude classique d'allocation de n boules dans C cellules [28]. Le cas qui nous intéresse est celui où les λ_i sont vus comme des variables indépendantes de distribution gamma de paramètres (α, β) .

Ainsi les p_i peuvent être représentés comme :

$$p_i = \Gamma(\alpha, \beta) (\Gamma(\alpha, \beta) + \Gamma(\alpha(C-1), \beta))^{-1} = \Gamma(\alpha, 1) (\Gamma(\alpha, 1) + \Gamma(\alpha(C-1), 1))^{-1},$$

où les variables $\Gamma(\alpha, 1)$ et $\Gamma(\alpha(C-1), 1)$ sont indépendantes.

Ainsi p_i ne dépend pas du paramètre β et possède une distribution Beta de paramètres $(\alpha, \alpha(C-1))$ [35]. Il en découle que le vecteur $(p_1, \dots, p_C)$ possède une distribution de Dirichlet de paramètres $\alpha_i = \alpha, i = 1, \dots, C$, et le vecteur $(n_1, \dots, n_C)$ une distribution de Dirichlet composée avec une distribution multinomiale [37], de la forme :

$$\mathbb{P}[n_1 = k_1, \dots, n_C = k_C] = \frac{n!}{k_1! \dots k_C!} \frac{\Gamma(\alpha C)}{\Gamma(\alpha)^C} \frac{\prod_{i=1}^C \Gamma(\alpha + k_i)}{\Gamma(\alpha C + n)},$$

qui ne dépend toujours pas de β .

On définit $m(n, C, \alpha, j) = \mathbb{E}[v(n, C, j)]$. Il est montré dans [10] que le tableau $\{m(n, C, \alpha, j), j = 0, \dots, n\}$ donne une bonne caractérisation du recrutement.

Pour étudier les propriétés de $m(n, C, \alpha, j)$, nous utilisons la représentation $v(n, C, j) = \sum_{i=1}^C \chi_i(n, C, j)$, où $\chi_i(n, C, j) = 1$ si le centre i a recruté exactement j patients, et 0 sinon. On a $\mathbb{P}[\chi_i(n, C, j) = 1] = \mathbb{P}[n_i = j]$. Ainsi :

$$\mathbb{E}[v(n, C, j)] = \sum_{i=1}^C \mathbb{P}[n_i = j]. \quad (1.3)$$

Lemme 1 Si les λ_i sont tous gamma distribuée de paramètres (α, β) , alors

$$m(n, C, \alpha, j) = C \binom{n}{j} \frac{\mathcal{B}(\alpha + j, \alpha(C-1) + n - j)}{\mathcal{B}(\alpha, \alpha(C-1))}, j = 0, \dots, n$$

Preuve En utilisant la représentation (1.3), et en utilisant le fait que pour tout i , p_i possède une distribution Beta et n_i une distribution Beta-binomiale, ainsi :

$$\mathbb{P}[n_i = j] = \mathbb{E}[\mathbb{P}(n_i = j | p_i)] = \binom{n}{j} \mathbb{E}[p_i^j (1 - p_i)^{n-j}],$$

en appliquant la formule (35.154) de [37], la formule est prouvée.

Une étude par simulation réalisée dans [\[10\]](#) permet d'observer ces courbes d'occupations et de valider le modèle.

Chapitre 2

Enrichissement du modèle

Depuis sa création de nombreux enrichissements ont été ajoutés au modèle Poisson-gamma, permettant ainsi d'augmenter le nombre d'applications du modèle sur des situations réelles.

Il s'agit surtout d'améliorations dont la source vient principalement de l'observation sur de multiples jeux de données de phénomènes se produisant pendant des dynamiques de recrutement, phénomènes qui soit remettent en question les hypothèses même du modèle, ou dans certaines situations, ajoutent des données complémentaires observables et qui peuvent être recueillies permettant ainsi d'enrichir dans ce cas le modèle.

2.1 Sensibilité du modèle aux paramètres de la dynamique de recrutement ¹

Le modèle Poisson-gamma peut être enrichi de multiples manières, le problème étant de ne pas apporter trop de complexité aux niveaux des calculs et des analyses des résultats. Depuis sa publication de nombreuses études ont été menées en ce sens pour l'étude de l'effet de la randomisation des centres stratifiés ou non [4], à la modélisation de la prédiction d'événement [5], et également à la prévision de la randomisation (et non plus seulement l'inclusion) par des processus et la prise de décision quant au recrutement et la fermeture de certains centres investigateurs [6].

Nous pouvons également de nouveau citer [11] portant sur les drop-out en cours de randomisation vu en section 2.3, ainsi que le modèle de coût [45] introduit au paragraphe 3.4.

Le modèle Poisson-gamma est ainsi de plus en plus utilisé mais deux problématiques émergent lors de son application sur données réelles :

- **Quelle est la performance du modèle :** Cette question a été étudiée dans [46] par une approche théorique. Les résultats sont essentiellement asymptotiques. Ainsi des bornes sur les valeurs des paramètres sont nécessaires pour assurer la validité du modèle.

1. Ce travail a fait l'objet d'une soumission : Proceedings of the 8th International Workshop on Simulation, Vienne, 2015. [48]

- **L'hypothèse des intensités constantes au cours du temps est-elle réaliste :** Le modèle peut dans tout les cas être enrichi pour prendre en compte les interruptions de recrutement ou les changements d'intensités mais l'apport de nouveaux paramètres explicatifs va de pair avec un ajout de sources d'erreurs de prédiction ainsi que des difficultés de collectes des données nécessaires à l'estimation des paramètres. Il est donc naturel de se demander quand est-ce que ces ajouts sont nécessaires non plus seulement en terme de performances prédictive.

Le but de cette section est, avant même de définir un modèle sur la dynamique de recrutement (quelque soit l'étude de la modélisation) qui utilise de manière sous-jacente la modélisation Poisson-gamma, d'étudier la sensibilité par simulation de ce dernier quant à certains paramètres intrinsèques à une dynamique de recrutement (nombre de centres, pauses dans le recrutement, dates d'ouvertures inconnues ou mal fournies, changement linéaire d'intensité au cours du temps) en terme de performance prédictive. Ceci permettant ainsi par la suite de pouvoir décider si le phénomène doit être pris en compte (trop forte erreur de prédiction) ou bien accepter l'erreur commise (erreur trop faible pour que les difficultés de collectes et de paramétrisations citées plus haut ne doivent être incorporées au modèle)

Pour ce faire quatre scénarios sont introduit, chacun étudiant particulièrement un des possibles ajout de biais de prédiction.

Scénario 1 : Sensibilité aux nombre de centres.

Scénario 2 : Sensibilité aux interruptions de recrutement.

Scénario 3 : Sensibilité aux dates d'ouvertures.

Scénario 4 : Sensibilité aux changements d'intensités aux cours du temps.

La méthode de génération des données nécessaires aux études de comparaison sont données par la suite selon les différents scénarios investigués.

2.1.1 Procédures de générations des données

Généralités

Considérons un essai multicentrique incluant C centres et dont le NSN est de $n = 1000$. La durée attendue est de 156 semaines (3 ans). La performance prédictive du modèle selon le scénario envisagé est étudiée par la biais d'un plan de simulation. La génération des données se fait en deux étapes :

- La génération de R processus de recrutement globaux $\{N^r(t), 0 \leq t \leq T^r, 1 \leq r \leq R\}$ où T^r dénotes le premier temps vérifiant $N^r(T^r) = 1000$. Ces processus sont générés selon des paramètres dépendant du scénario.
- Sachant t_1 donné :
 1. Les paramètres (α, μ) du modèle Poisson-gamma sont estimés par application du Théorème 1.
 2. La durée moyenne de l'essai $T_{t_1}^r$ est calculée par l'application du Théorème 3 à l'instant t_1 .

3. La performance prédictive du modèle est mesurée en terme d'erreur relative par la formule :

$$E_{t_1}^r = \frac{T_{t_1}^r - T^r}{T^r}. \quad (2.1)$$

Génération de la dynamique de recrutement : Scénario 1

Ce scénario permet l'étude de sensibilité du modèle quant aux nombres de centres inclus dans l'étude. Des résultats théoriques sont déjà disponibles (supérieur à 50) et sur données réelles (supérieur à 20), cette étude rajoute des connaissances par simulation quant à cette problématique. Forcément le résultat est très dépendant des taux de recrutement propres aux centres, c'est pourquoi les deux paramètres sont étudiés en même temps. Les variations entre deux générations portent sur la valeur de C ainsi que sur la valeur de α (β est fixé à 2). Pour $C = 1, 2, \dots, 25$ et $\alpha = 0.7, 0.8, \dots, 3.5$ sont effectuées 250 générations de l'algorithme suivant :

1. Génération des taux de recrutement selon une loi $Gamma(\alpha, \beta)$ pour chaque centre.
2. Génération de la trajectoire de recrutement pour chacun des C centres.
3. Agrégation des C trajectoires pour obtenir la trajectoire globale.
4. Identification du temps nécessaire pour recruter $n = 1000$ patients.
5. Calcul de l'erreur relative aux instants intermédiaires $t_1 = 78$ semaines et $t_1 = 104$ semaines (1/2 et 2/3 des 156 semaines attendues).

Génération de la dynamique de recrutement : Scénario 2

Ce scénario étudie l'impact des interruptions de recrutement au cours de la phase de recrutement. Ainsi l'erreur de prédiction sera mise en relation avec la durée d'inactivité totale des centres pendant le recrutement.

C est ici fixé à 20, et les valeurs α et β sont respectivement fixées à 0.7 et 2. Notons par la suite π_d^b la proportion d'interruptions de durée d et γ l'intensité des interruptions. Les valeurs de d considérées sont 0, 2, 4, 8, 12, et 24 semaines et les valeurs de γ sont 1, 2 et 3. Les configurations choisies (i.e. les vecteurs $\vec{\pi}^b = (\pi_0^b, \pi_2^b, \pi_4^b, \pi_8^b, \pi_{12}^b, \pi_{24}^b)$) sont données dans le Tableau 2.1.

Pour chaque centre le taux de recrutement ne sera pas constant, son aspect est visible en Figure 2.1 dont certaines notations sont définies ci-dessous.

La génération se fait comme suit :

1. Génération des intensités selon une loi $Gamma(\alpha, \beta)$.
2. On fixe l'intensité des breaks γ et une configuration $\vec{\pi}^b$.
3. On effectue 100 générations selon l'algorithme :
 - (a) Génération des instant d'interruptions B_- pour chaque configuration et pour chaque centre selon une loi exponentielle de paramètres γ .
 - (b) Génération des durées d'interruptions D pour chaque centre par l'utilisation d'une loi multinomiale de paramètres $\vec{\pi}^b$.
 - (c) Génération du prochain instant d'interruptions $B_+ = B_- + D + E$ où E est généré selon une loi exponentielle de paramètre γ .

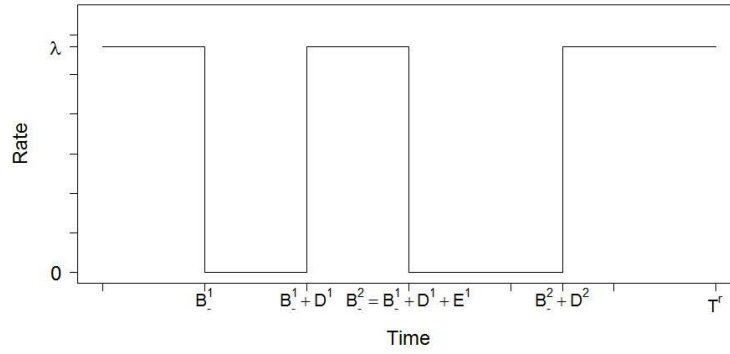


FIGURE 2.1 – Aspect des intensités de recrutement d'un centre pour le Scénario 2.

- (d) Génération des dynamiques de recrutement entre $B_- + D$ et B_+ pour chaque centre.
 - (e) Retour à a) jusqu'à ce que chaque centre recrute 1000 patients en remplaçant B_- par B_+ .
 - (f) Agrégation des C dynamiques de recrutement pour obtenir la dynamique globale.
 - (g) Raccourcissement de la trajectoire globale pour atteindre $n = 1000$ patients.
 - (h) Identification du temps T^r et calcul de $DI^{b,r}$ qui est la moyenne, sur les centres, de la durée d'inactivité et $PI^{b,r} = \frac{DI^r}{T^r}$ la proportion moyenne de temps inactif.
 - (i) Calcul de l'erreur relative aux instants intermédiaires $t_1 = 78$ semaines et $t_1 = 104$ semaines.
4. Calcul de l'erreur relative moyenne aux temps t_1 , de DI^b la moyenne des $DI^{b,r}$'s et PI^b la moyenne des $PI^{b,r}$'s.

Génération de la dynamique de recrutement : Scénario 3

Ce scénario a pour but de souligner l'impact des délais sur le recrutement (les dates d'ouvertures non spécifiées).

Ainsi on considère C fixé à 20, $\alpha = 0.7$ et $\beta = 2$.

Désignons par π_d^o la proportion de centres dont les dates d'ouvertures sont augmentées de d semaines. Les valeurs de d considérées sont 0, 1, 2, 4, 8, 12, 16, 20 et 24 semaines. Les configurations choisies (i.e. les vecteurs $\vec{\pi}^o = (\pi_0^o, \pi_1^o, \pi_2^o, \pi_4^o, \pi_8^o, \pi_{12}^o, \pi_{16}^o, \pi_{20}^o, \pi_{24}^o)$) sont données dans le Tableau 2.2.

Pour chaque centre le taux de recrutement n'est de nouveau pas constant, son aspect est visible sur la Figure 2.2 dont les notations sont définies ci-dessous.

1. Génération des intensités selon une loi $Gamma(\alpha, \beta)$.
2. On fixe une configuration $\vec{\pi}^o$.
3. On effectue 100 générations selon l'algorithme :

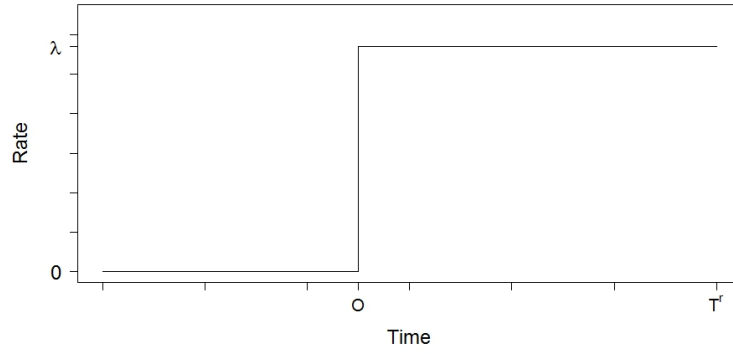


FIGURE 2.2 – Aspect des intensités de recrutement d'un centre pour le Scénario 3 .

- (a) Génération des dates d'ouvertures (vu comme des délais non connus du modèle) O_i pour chaque centre i en utilisant une loi multinomiale de paramètre $\vec{\pi}^o$.
 - (b) Génération des dynamiques de recrutement à partir de O_i pour chaque centre i jusqu'à ce que $n = 1000$ patients soient atteints.
 - (c) Agrégation des C dynamiques de recrutement.
 - (d) Raccourcissement de la trajectoire globale pour atteindre $n = 1000$ patients.
 - (e) Identification du temps T^r et calcul de $DI^{o,r}$ qui est la moyenne, sur les centres, de la durée de délai et $PI^{o,r} = \frac{DI^{o,r}}{T^r}$ la proportion moyenne de temps inactif.
 - (f) Calcul de l'erreur relative aux instants intermédiaires $t_1 = 78$ semaines et $t_1 = 104$ semaines.
4. Calcul de l'erreur relative moyenne aux temps t_1 , de DI^o la moyenne des $DI^{o,r}$'s et PI^o la moyenne des $PI^{o,r}$'s.

Génération de la dynamique de recrutement : Scénario 4

Le but de ce scénario est d'étudier l'impact de changements continus des intensités de recrutement au cours du temps. C fixé à 20, $\alpha = 0.7$ et $\beta = 2$.

Sont considérées deux modifications des intensités :

- Une période d'apprentissage : les intensités croissent linéairement de 0 jusqu'à leur valeur maximum prédéfinie. Notons π_d^l la proportion de centres dont la phase d'apprentissage est de d semaines. Les valeurs de d considérées sont 0, 2, 4 et 8 semaines. Les configurations choisies (i.e. les vecteurs $\vec{\pi}^l = (\pi_0^l, \pi_2^l, \pi_4^l, \pi_8^l)$) sont données dans le Tableau 2.3.
- Une période de tarissement : à partir d'un instant T^{du} , l'intensité décroît linéairement avec une pente définie par la suite. Soit π_k^{du} la proportion de centres dont la pente de décroissance de l'intensité est s_k . Les valeurs de s_k considérées sont $s_1 = 0$, $s_2 = -0.05$, $s_3 = -0.1$ et $s_4 = -0.2$. Les valeurs du début de la décroissance T^{du} considérées sont 108 et 120 semaines. Les configurations choisies (i.e. les vecteurs $\vec{\pi}^{du} = (\pi_1^{du}, \pi_2^{du}, \pi_3^{du}, \pi_4^{du})$) sont données au Tableau 2.3.

Pour chaque centre les intensités de recrutement ne sont pas constantes et leur aspect est donné en Figure 2.3 dont les notations sont définies ci-dessous.

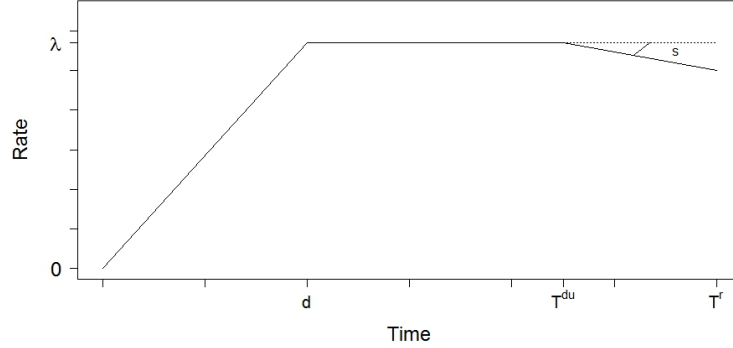


FIGURE 2.3 – Aspect des intensités de recrutement d'un centre pour le Scénario 4.

1. Génération des intensités selon une loi $Gamma(\alpha, \beta)$.
2. On fixe une configuration $\vec{\pi}^l$ pour la période d'apprentissage et une configuration $\vec{\pi}^{du}$ pour la période de tarissement.
3. On effectue 100 générations selon l'algorithme :
 - (a) Génération de la fin de la période d'apprentissage pour chaque centre selon une loi multinomiale de paramètre $\vec{\pi}^l$.
 - (b) Génération de la pente de la période de tarissement pour chaque centre selon une loi multinomiale de paramètre $\vec{\pi}^{du}$.
 - (c) Génération des dynamiques de recrutement à partir de O_i pour chaque centre i jusqu'à ce que $n = 1000$ patients soient atteints.
 - (d) Agrégation des C dynamique de recrutement.
 - (e) Raccourcissement de la trajectoire globale pour atteindre $n = 1000$ patients.
 - (f) Identification de la durée de recrutement et calcul du ratio :

$$F^r = 1 - \frac{\sum_{i=1}^C \mathcal{A}_i(T^r)}{\sum_{i=1}^C \lambda_i T^r}, \quad (2.2)$$

où $\mathcal{A}_i(T^r)$ est l'aire sous la courbe de l'intensité du centre i entre 0 et T^r comme visible sur la Figure 2.3. F^r représente la perte relative d'intensité entre une situation perturbée selon les configurations envisagées et une intensité constante.

- (g) Calcul de l'erreur relative aux instants intermédiaires $t_1 = 78$ semaines et $t_1 = 104$ semaines.

4. Calcul de l'erreur relative moyenne et F la moyenne des F^r 's.

Remarque Notons que le scénario 1 diffère des 3 autres scénarios dans le sens où les données générées ne sont pas modifiées par des configurations. Seule la précision de la prédiction (sans ajout de mauvaises spécifications) est évaluée. Le paramètre d'intérêt est une moyenne mais la multiplicité des jeux de données générés nous permet d'évaluer le comportement du modèle au travers de la variabilité de ces données.

2.1.2 Résultats

Sensibilité de la prévision de la durée de recrutement selon le nombre de centres et la moyenne de recrutement (scénario 1)

La Figure 2.4 est le graphique représentant l'erreur relative commise de prédiction (un point blanc signifie une erreur relative inférieure à 10% et inversement pour les points noirs) comme une fonction de C en abscisse et de α en ordonnée.

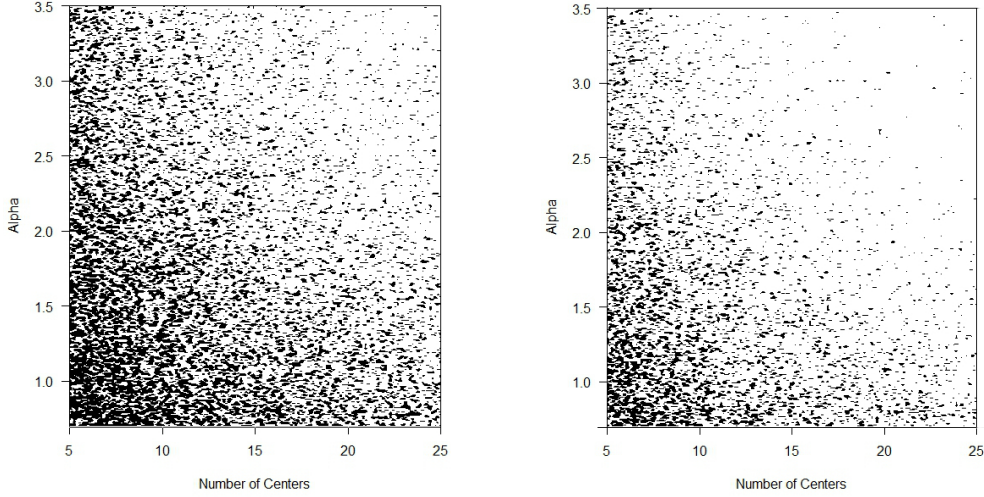


FIGURE 2.4 – Erreur relative de prédiction de la durée de l'essai des données générées jusqu'à $\theta=78$ semaines (a) et $\theta=104$ semaines (b).

Plusieurs conclusions découlent de l'étude de la Figure 2.4. On constate premièrement que comme attendu le paramètre C (respectivement α) a un impact significatif sur la précision de la prédiction, en effet moins le nombre de centres sera élevé, plus la prédiction sera en moyenne supérieure à 10% d'erreur relative.

On constate deuxièmement un phénomène d'atténuation (respectivement d'ajout) entre les deux paramètres C et α . Si un des deux est petit, la précision peut être importante si l'autre paramètre est élevé (bord gauche en haut et bord droit en bas). Cependant si les deux sont petits alors presque plus aucune prévision n'est proche relativement à la durée réelle.

Finalement ces phénomènes sont amplifiés (respectivement atténués) lorsque $t_1 = 78$ semaines (resp. $t_1 = 104$ semaines).

Sensibilité de la prévision de la durée de recrutement selon les interruptions de recrutement (scénario 2)

Les résultats des simulations aux temps $\theta = 78$ semaines et $\theta = 104$ semaines sont collectés dans le Tableau 2.1.

Dans l'optique de souligner l'impact de chacun des paramètres de ce scénario sur l'erreur relative, une ANOVA à deux facteurs est appliquée. Le premier facteur est γ et le second $\bar{\pi}^b$. Les modalités de ces facteurs sont données par les niveaux 1 to 18 du Tableau 2.1 (les niveaux 19, 20 et 21 étudient le cas où toutes les probabilités sont égales de sorte que l'analyse soit complète quant à l'effet du paramètre γ). Il n'y a pas

TABLEAU 2.1 – Resultats du Scénario 2 : les différentes valeurs γ et π^b des paramètres, la caractéristique PI^b de chaque configuration et les resultat en terme d’erreur relative selon t_1 . Les valeurs sont données en pourcentage.

Niveau	Valeurs des paramètres		PI^b	Erreur relative moyenne à $t_1 =$	
	γ	π^b		78 semaines	104 semaines
1	1	$(\frac{1}{2}, \frac{1}{10}, \frac{1}{10}, \frac{1}{10}, \frac{1}{10}, \frac{1}{10})$	3.57	2.88(± 0.10)	1.45(± 0.55)
2	$\frac{1}{2}$	$(\frac{1}{2}, \frac{1}{10}, \frac{1}{10}, \frac{1}{10}, \frac{1}{10}, \frac{1}{10})$	7.30	6.57	2.99
3	$\frac{1}{3}$	$(\frac{1}{2}, \frac{1}{10}, \frac{1}{10}, \frac{1}{10}, \frac{1}{10}, \frac{1}{10})$	10.90	9.27	4.51
4	1	$(\frac{1}{10}, \frac{1}{2}, \frac{1}{10}, \frac{1}{10}, \frac{1}{10}, \frac{1}{10})$	4.01	3.79	1.81
5	$\frac{1}{2}$	$(\frac{1}{10}, \frac{1}{2}, \frac{1}{10}, \frac{1}{10}, \frac{1}{10}, \frac{1}{10})$	8.17	7.53	3.43
6	$\frac{1}{3}$	$(\frac{1}{10}, \frac{1}{2}, \frac{1}{10}, \frac{1}{10}, \frac{1}{10}, \frac{1}{10})$	11.79	10.82	5.65
7	1	$(\frac{1}{10}, \frac{1}{10}, \frac{1}{2}, \frac{1}{10}, \frac{1}{10}, \frac{1}{10})$	4.65	4.09	2.01
8	$\frac{1}{2}$	$(\frac{1}{10}, \frac{1}{10}, \frac{1}{2}, \frac{1}{10}, \frac{1}{10}, \frac{1}{10})$	9.24	9.04	3.93
9	$\frac{1}{3}$	$(\frac{1}{10}, \frac{1}{10}, \frac{1}{2}, \frac{1}{10}, \frac{1}{10}, \frac{1}{10})$	13.97	13.64	6.46
10	1	$(\frac{1}{10}, \frac{1}{10}, \frac{1}{10}, \frac{1}{2}, \frac{1}{10}, \frac{1}{10})$	5.70	5.30	2.38
11	$\frac{1}{2}$	$(\frac{1}{10}, \frac{1}{10}, \frac{1}{10}, \frac{1}{2}, \frac{1}{10}, \frac{1}{10})$	11.31	10.86	5.32
12	$\frac{1}{3}$	$(\frac{1}{10}, \frac{1}{10}, \frac{1}{10}, \frac{1}{2}, \frac{1}{10}, \frac{1}{10})$	17.43	17.78	8.43
13	1	$(\frac{1}{10}, \frac{1}{10}, \frac{1}{10}, \frac{1}{10}, \frac{1}{2}, \frac{1}{10})$	6.77	6.59	3.08
14	$\frac{1}{2}$	$(\frac{1}{10}, \frac{1}{10}, \frac{1}{10}, \frac{1}{10}, \frac{1}{2}, \frac{1}{10})$	13.59	13.75	6.13
15	$\frac{1}{3}$	$(\frac{1}{10}, \frac{1}{10}, \frac{1}{10}, \frac{1}{10}, \frac{1}{2}, \frac{1}{10})$	20.12	21.91	10.01
16	1	$(\frac{1}{10}, \frac{1}{10}, \frac{1}{10}, \frac{1}{10}, \frac{1}{10}, \frac{1}{2})$	9.96	8.65	4.23
17	$\frac{1}{2}$	$(\frac{1}{10}, \frac{1}{10}, \frac{1}{10}, \frac{1}{10}, \frac{1}{10}, \frac{1}{2})$	20.29	19.23	9.22
18	$\frac{1}{3}$	$(\frac{1}{10}, \frac{1}{10}, \frac{1}{10}, \frac{1}{10}, \frac{1}{10}, \frac{1}{2})$	30.21	32.53	13.98
19	1	$(\frac{1}{6}, \frac{1}{6}, \frac{1}{6}, \frac{1}{6}, \frac{1}{6}, \frac{1}{6})$	5.81	5.26	2.53
20	$\frac{1}{2}$	$(\frac{1}{6}, \frac{1}{6}, \frac{1}{6}, \frac{1}{6}, \frac{1}{6}, \frac{1}{6})$	11.76	11.53	4.97
21	$\frac{1}{3}$	$(\frac{1}{6}, \frac{1}{6}, \frac{1}{6}, \frac{1}{6}, \frac{1}{6}, \frac{1}{6})$	17.18	17.52	7.96

d’interaction significative entre ces deux facteurs, cependant chaque facteur a un effet significatif sur l’erreur relative.

De plus, étant donné une configuration, l’erreur relative moyenne croît avec les valeurs de γ et sachant γ donné elle croît avec la valeur maximale de la plus grande probabilité de la configuration. L’étude est finalisée par une application d’un Test de Tuckey [62]. Ce test permet d’identifier les niveaux des facteurs qui ne sont pas significativement différents en terme d’erreur relative (voir la Figure 2.5), et ainsi de définir un seuil sous lequel il n’est pas nécessaire de considérer la perturbation selon une erreur maximum acceptée.

L’interprétation de la Figure 2.5 n’est pas triviale car elle est reliée aux différents niveaux (paramètre composite). Il est plus pratique de considérer PI^b . La Figure 2.6 représente l’erreur relative moyenne comme fonction de PI^b . Sachant un degré d’erreur

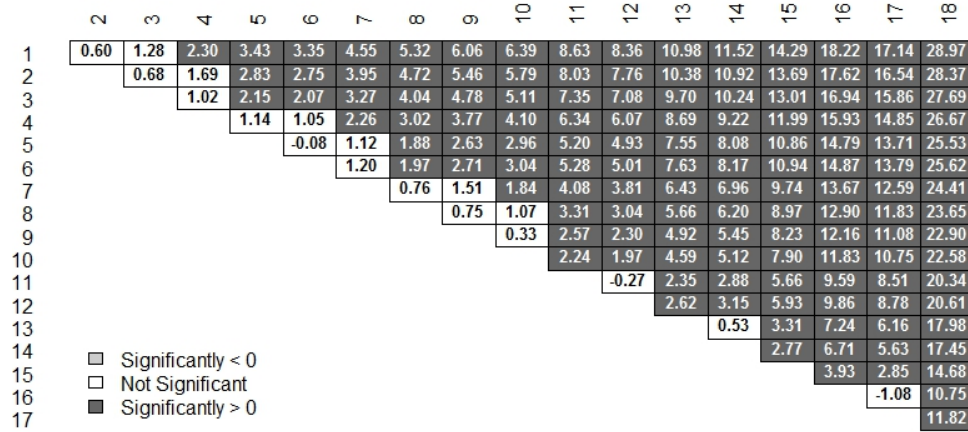


FIGURE 2.5 – Tests de Tukey's HSD (honest significant difference) aux instants $t_1=104$ semaines. Les valeurs à la colonne i et la ligne j est le fonction discriminante d'un test de comparaison entre les moyennes des niveaux i et j (ordonnés selon leur valeur de PI^b) basé sur l'écart d'une distribution studentisée (méthode HSD). Cette méthode prend en compte l'inflation du risque α du fait de la multiplicité des tests de comparaison.

acceptable défini, il est facile de trouver les niveaux acceptables.

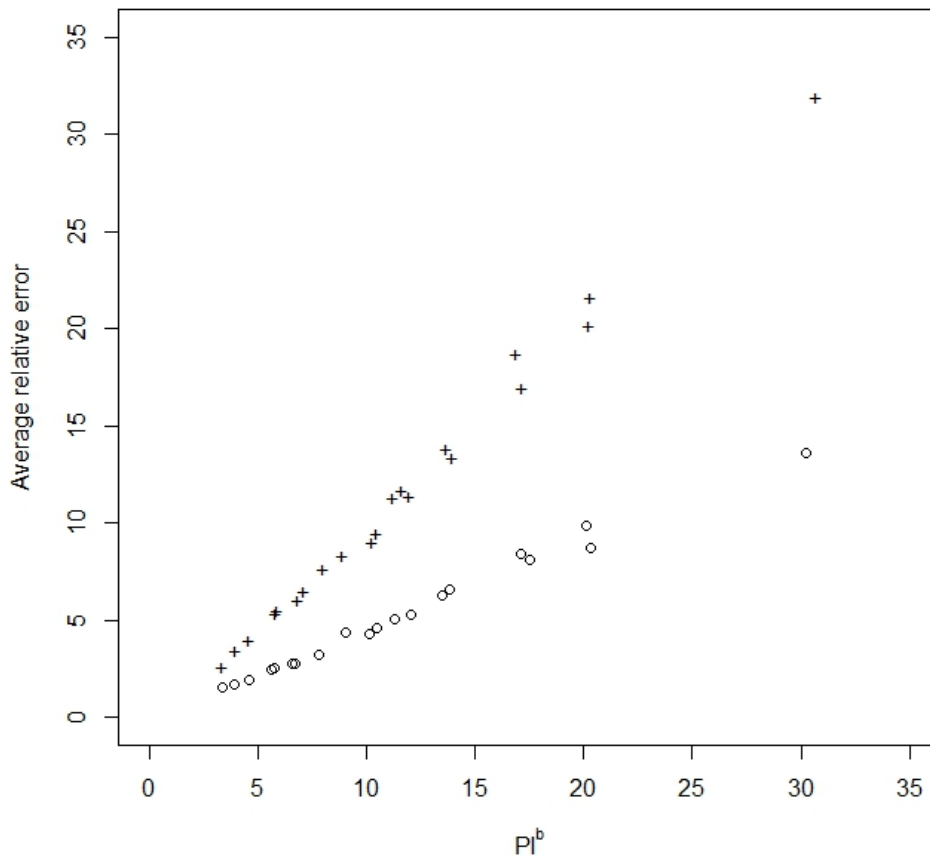


FIGURE 2.6 – Erreur relative moyenne de la prédiction de la durée de recrutement à $t_1 = 78$ semaines, '+' et $t_1 = 104$ semaines, 'o' comme fonction de PI^b .

Sensibilité de la prévision de la durée de recrutement selon les délais des dates d'ouvertures (scénario 3)

Les résultats de simulation sont collectés dans le Tableau 2.2.

TABLEAU 2.2 – Résultats du Scénario 3 : les différentes valeurs de $\vec{\pi}^o$ des paramètres, la caractéristique PI^o de chaque configuration et les résultats en terme d'erreur relative moyenne aux différents temps t_1 . Les valeurs sont exprimées en pourcentage.

Niveau	Valeurs des paramètres	PI^o	Temps intermédiaires $t_1 =$	
	$\vec{\pi}^o$		78 semaines	104 semaines
1	$(\frac{1}{9}, \frac{1}{9}, \frac{1}{9}, \frac{1}{9}, \frac{1}{9}, \frac{1}{9}, \frac{1}{9}, \frac{1}{9}, \frac{1}{9})$	6.67	7.64(± 0.77)	3.94(± 0.38)
2	$(\frac{1}{2}, \frac{1}{16}, \frac{1}{16}, \frac{1}{16}, \frac{1}{16}, \frac{1}{16}, \frac{1}{16}, \frac{1}{16}, \frac{1}{16})$	4.57	4.30	2.11
3	$(\frac{1}{16}, \frac{1}{2}, \frac{1}{16}, \frac{1}{16}, \frac{1}{16}, \frac{1}{16}, \frac{1}{16}, \frac{1}{16}, \frac{1}{16})$	4.79	4.13	2.21
4	$(\frac{1}{16}, \frac{1}{16}, \frac{1}{2}, \frac{1}{16}, \frac{1}{16}, \frac{1}{16}, \frac{1}{16}, \frac{1}{16}, \frac{1}{16})$	5.07	4.73	2.49
5	$(\frac{1}{16}, \frac{1}{16}, \frac{1}{16}, \frac{1}{2}, \frac{1}{16}, \frac{1}{16}, \frac{1}{16}, \frac{1}{16}, \frac{1}{16})$	5.16	5.88	2.52
6	$(\frac{1}{16}, \frac{1}{16}, \frac{1}{16}, \frac{1}{16}, \frac{1}{2}, \frac{1}{16}, \frac{1}{16}, \frac{1}{16}, \frac{1}{16})$	6.28	7.27	3.39
7	$(\frac{1}{16}, \frac{1}{16}, \frac{1}{16}, \frac{1}{16}, \frac{1}{16}, \frac{1}{2}, \frac{1}{16}, \frac{1}{16}, \frac{1}{16})$	7.12	8.74	4.06
8	$(\frac{1}{16}, \frac{1}{16}, \frac{1}{16}, \frac{1}{16}, \frac{1}{16}, \frac{1}{16}, \frac{1}{2}, \frac{1}{16}, \frac{1}{16})$	8.02	10.74	7.55
9	$(\frac{1}{16}, \frac{1}{16}, \frac{1}{16}, \frac{1}{16}, \frac{1}{16}, \frac{1}{16}, \frac{1}{16}, \frac{1}{2}, \frac{1}{16})$	9.10	12.53	6.03
10	$(\frac{1}{16}, \frac{1}{16}, \frac{1}{16}, \frac{1}{16}, \frac{1}{16}, \frac{1}{16}, \frac{1}{16}, \frac{1}{16}, \frac{1}{2})$	9.92	14.26	6.42

Une analyse de variance à un facteur permet de conclure à un effet significatif du facteur $\vec{\pi}^o$, à savoir des délais dans les dates d'ouvertures du recrutement, sur l'erreur relative.

L'interprétation est plus facile car ne comprenant qu'un seul facteur mais la caractéristique PI^o permet comme pour le scénario précédent de trouver facilement le seuil d'acceptabilité de la perturbation.

En effet il est facilement observable sur la Figure 2.7 que l'erreur relative moyenne de prédiction augmente avec la moyenne des proportion de temps d'indicativité PI^o .

Sensibilité de la prévision de la durée de recrutement selon les changements d'intensités continus au cours du temps (scénario 4)

Les résultats de simulation sont collectés dans le Tableau 2.3.

Une analyse de variance à 3 facteurs est effectuée pour étudier les effets de chacun des facteurs et des interactions possibles. Le premier facteur est $\vec{\pi}^l$, le second T^{du} et le troisième est $\vec{\pi}^{du}$. Les modalités sont données par les niveaux 1 à 50 du Tableaux 2.3. De nouveau aucune interactions entre facteur n'est significative mais chaque facteur a un effet significatif sur l'erreur relative de prédiction. De nouveau l'interprétation est plus aisée lorsque l'on considère la caractéristique $\mathbf{F}$ et les mêmes conclusions sont observées pour la définition d'un seuil d'acceptabilité de la perturbation visible sur la Figure 2.8.

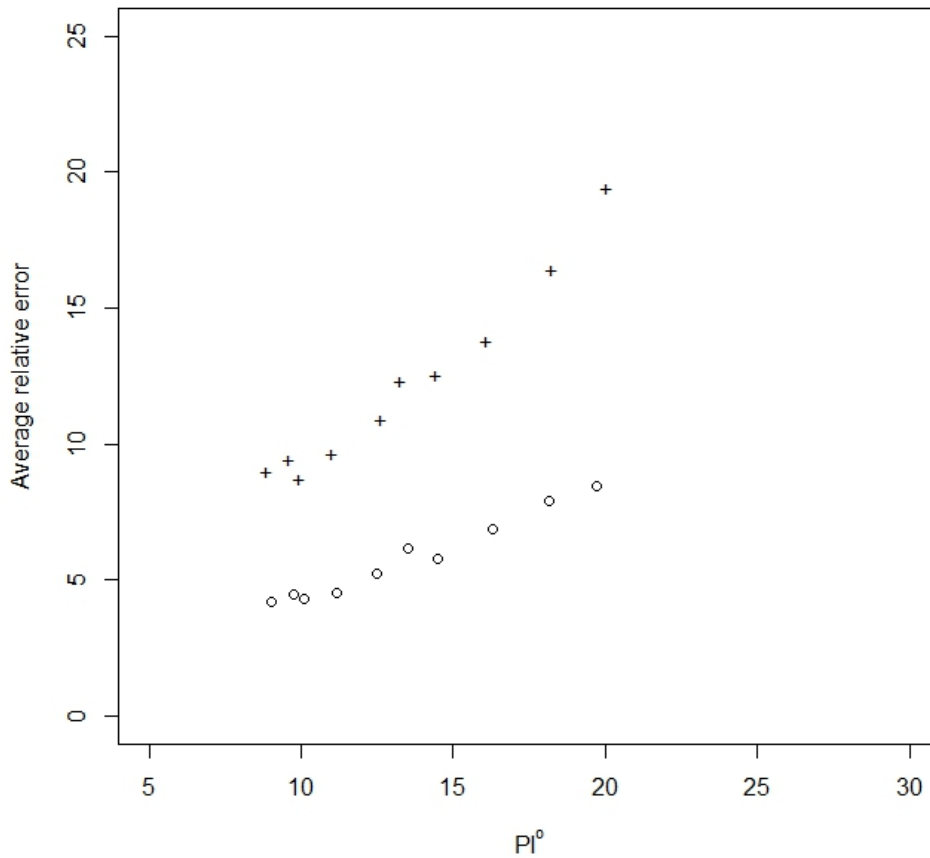


FIGURE 2.7 – Erreur relative moyenne de la prédiction de la durée de recrutement à $t_1 = 78$ semaines, '+' et $t_1 = 104$ semaines, 'o' comme fonction de PI^0 .

Remarque Notons que le phénomène de la phase de tarissement amène à une sous-estimation de la durée de recrutement au contraire de toutes les autres perturbations.

2.1.3 Recommandations

Plusieurs conclusions émanent de ces simulations :

- Premièrement l'importance du temps intermédiaire t_1 choisi pour estimer les paramètres sur l'erreur relative de la prédiction de la durée du recrutement. En fait, plus l'étude intermédiaire se fait tôt, plus grand sera l'impact de la perturbation sur l'erreur de prédiction commise.
- Deuxièmement le modèle continue tout de même à donner des résultats convaincant même pour un nombre de centres inférieur à 20 tant que les intensités de recrutement sont assez élevées (voir Figure 2.4). Il n'existe pas actuellement de relation explicite entre C , α et l'erreur relative de prédiction mais des études par simulations *ad hoc* permettent d'assurer la validité de l'application du modèle pour des valeurs de paramètres spécifiques.
- Troisièmement les différentes perturbations introduites dans le modèle jouent un rôle dans son efficacité. De nouveau aucune relation explicite n'existe entre

TABLEAU 2.3 – Résultats du Scénario 4 : les différentes valeurs de $\bar{\pi}^l$, T_d et $\bar{\pi}^{du}$ des paramètres, la caractéristique **F** de chaque configuration et les résultats en terme d'erreur relative moyenne aux différents temps t_1 . Les valeurs sont exprimées en pourcentage sauf T_d qui est en semaines.

Niveau	Valeurs des paramètres			F	Temps intermédiaires $t_1 =$	
	$\bar{\pi}^l$	T_d	$\bar{\pi}^{du}$		78 semaines	104 semaines
1	$(\frac{1}{4}, \frac{1}{4}, \frac{1}{4}, \frac{1}{4})$	130	$(\frac{1}{2}, \frac{1}{6}, \frac{1}{6}, \frac{1}{6})$	6.05	-7.23(± 0.75)	-7.38(± 0.65)
2	$(\frac{1}{4}, \frac{1}{4}, \frac{1}{4}, \frac{1}{4})$	130	$(\frac{1}{6}, \frac{1}{2}, \frac{1}{6}, \frac{1}{6})$	8.46	-10.72	-11.35
3	$(\frac{1}{4}, \frac{1}{4}, \frac{1}{4}, \frac{1}{4})$	130	$(\frac{1}{6}, \frac{1}{6}, \frac{1}{2}, \frac{1}{6})$	9.70	-12.24	-13.06
4	$(\frac{1}{4}, \frac{1}{4}, \frac{1}{4}, \frac{1}{4})$	130	$(\frac{1}{6}, \frac{1}{6}, \frac{1}{6}, \frac{1}{2})$	10.98	-14.26	-14.91
5	$(\frac{1}{2}, \frac{1}{6}, \frac{1}{6}, \frac{1}{6})$	130	$(\frac{1}{4}, \frac{1}{4}, \frac{1}{4}, \frac{1}{4})$	8.33	-11.77	-11.77
6	$(\frac{1}{6}, \frac{1}{2}, \frac{1}{6}, \frac{1}{6})$	130	$(\frac{1}{4}, \frac{1}{4}, \frac{1}{4}, \frac{1}{4})$	8.37	-11.13	-11.40
7	$(\frac{1}{6}, \frac{1}{6}, \frac{1}{2}, \frac{1}{6})$	130	$(\frac{1}{4}, \frac{1}{4}, \frac{1}{4}, \frac{1}{4})$	9.05	-11.12	-11.59
8	$(\frac{1}{6}, \frac{1}{6}, \frac{1}{6}, \frac{1}{2})$	130	$(\frac{1}{4}, \frac{1}{4}, \frac{1}{4}, \frac{1}{4})$	9.54	-10.80	-11.31
9	$(\frac{1}{2}, \frac{1}{6}, \frac{1}{6}, \frac{1}{6})$	130	$(\frac{1}{2}, \frac{1}{6}, \frac{1}{6}, \frac{1}{6})$	5.50	-7.27	-7.44
10	$(\frac{1}{2}, \frac{1}{6}, \frac{1}{6}, \frac{1}{6})$	130	$(\frac{1}{6}, \frac{1}{2}, \frac{1}{6}, \frac{1}{6})$	8.19	-11.15	-11.35
11	$(\frac{1}{2}, \frac{1}{6}, \frac{1}{6}, \frac{1}{6})$	130	$(\frac{1}{6}, \frac{1}{6}, \frac{1}{2}, \frac{1}{6})$	9.43	-12.69	-13.66
12	$(\frac{1}{2}, \frac{1}{6}, \frac{1}{6}, \frac{1}{6})$	130	$(\frac{1}{6}, \frac{1}{6}, \frac{1}{6}, \frac{1}{2})$	10.15	-14.02	-14.69
13	$(\frac{1}{6}, \frac{1}{2}, \frac{1}{6}, \frac{1}{6})$	130	$(\frac{1}{2}, \frac{1}{6}, \frac{1}{6}, \frac{1}{6})$	5.95	-7.77	-7.27
14	$(\frac{1}{6}, \frac{1}{2}, \frac{1}{6}, \frac{1}{6})$	130	$(\frac{1}{6}, \frac{1}{2}, \frac{1}{6}, \frac{1}{6})$	8.19	-10.81	-11.28
15	$(\frac{1}{6}, \frac{1}{2}, \frac{1}{6}, \frac{1}{6})$	130	$(\frac{1}{6}, \frac{1}{6}, \frac{1}{2}, \frac{1}{6})$	9.67	-12.98	-13.56
16	$(\frac{1}{6}, \frac{1}{2}, \frac{1}{6}, \frac{1}{6})$	130	$(\frac{1}{6}, \frac{1}{6}, \frac{1}{6}, \frac{1}{2})$	10.46	-13.97	-14.47
17	$(\frac{1}{6}, \frac{1}{6}, \frac{1}{2}, \frac{1}{6})$	130	$(\frac{1}{2}, \frac{1}{6}, \frac{1}{6}, \frac{1}{6})$	5.91	-6.84	-7.78
18	$(\frac{1}{6}, \frac{1}{6}, \frac{1}{2}, \frac{1}{6})$	130	$(\frac{1}{6}, \frac{1}{2}, \frac{1}{6}, \frac{1}{6})$	8.63	-10.66	-10.96
19	$(\frac{1}{6}, \frac{1}{6}, \frac{1}{2}, \frac{1}{6})$	130	$(\frac{1}{6}, \frac{1}{6}, \frac{1}{2}, \frac{1}{6})$	9.91	-12.61	-13.42
20	$(\frac{1}{6}, \frac{1}{6}, \frac{1}{2}, \frac{1}{6})$	130	$(\frac{1}{6}, \frac{1}{6}, \frac{1}{6}, \frac{1}{2})$	10.79	-14.04	-14.47
21	$(\frac{1}{6}, \frac{1}{6}, \frac{1}{6}, \frac{1}{2})$	130	$(\frac{1}{2}, \frac{1}{6}, \frac{1}{6}, \frac{1}{6})$	6.36	-6.40	-7.27
22	$(\frac{1}{6}, \frac{1}{6}, \frac{1}{6}, \frac{1}{2})$	130	$(\frac{1}{6}, \frac{1}{2}, \frac{1}{6}, \frac{1}{6})$	9.08	-9.90	-11.11
23	$(\frac{1}{6}, \frac{1}{6}, \frac{1}{6}, \frac{1}{2})$	130	$(\frac{1}{6}, \frac{1}{6}, \frac{1}{2}, \frac{1}{6})$	10.36	-11.68	-13.12
24	$(\frac{1}{6}, \frac{1}{6}, \frac{1}{6}, \frac{1}{2})$	130	$(\frac{1}{6}, \frac{1}{6}, \frac{1}{6}, \frac{1}{2})$	10.99	-13.04	-14.41

ces paramètres et l'erreur relative de prédiction mais des études de simulations permettent de définir des seuils d'acceptabilité sous lesquels les perturbations peuvent être négligées (voir Figure 2.6, 2.7 et 2.8). C'est d'un intérêt tout parti-

Niveau	Valeurs des paramètres			F	Temps intermédiaires $t_1 =$	
	$\vec{\pi}^l$	T_d	$\vec{\pi}^{du}$		78 semaines	104 semaines
25	$(\frac{1}{4}, \frac{1}{4}, \frac{1}{4}, \frac{1}{4})$	143	$(\frac{1}{2}, \frac{1}{6}, \frac{1}{6}, \frac{1}{6})$	2.56	-1.31	-2.59
26	$(\frac{1}{4}, \frac{1}{4}, \frac{1}{4}, \frac{1}{4})$	143	$(\frac{1}{6}, \frac{1}{2}, \frac{1}{6}, \frac{1}{6})$	3.50	-3.61	-3.36
27	$(\frac{1}{4}, \frac{1}{4}, \frac{1}{4}, \frac{1}{4})$	143	$(\frac{1}{6}, \frac{1}{6}, \frac{1}{2}, \frac{1}{6})$	3.85	-4.23	-4.45
28	$(\frac{1}{4}, \frac{1}{4}, \frac{1}{4}, \frac{1}{4})$	143	$(\frac{1}{6}, \frac{1}{6}, \frac{1}{6}, \frac{1}{2})$	4.56	-4.79	-5.76
29	$(\frac{1}{2}, \frac{1}{6}, \frac{1}{6}, \frac{1}{6})$	143	$(\frac{1}{4}, \frac{1}{4}, \frac{1}{4}, \frac{1}{4})$	3.29	-4.15	-4.32
30	$(\frac{1}{6}, \frac{1}{2}, \frac{1}{6}, \frac{1}{6})$	143	$(\frac{1}{4}, \frac{1}{4}, \frac{1}{4}, \frac{1}{4})$	3.54	-3.51	-3.76
31	$(\frac{1}{6}, \frac{1}{6}, \frac{1}{2}, \frac{1}{6})$	143	$(\frac{1}{4}, \frac{1}{4}, \frac{1}{4}, \frac{1}{4})$	3.68	-3.26	-4.29
32	$(\frac{1}{6}, \frac{1}{6}, \frac{1}{6}, \frac{1}{2})$	143	$(\frac{1}{4}, \frac{1}{4}, \frac{1}{4}, \frac{1}{4})$	4.23	-3.16	-3.82
33	$(\frac{1}{2}, \frac{1}{6}, \frac{1}{6}, \frac{1}{6})$	143	$(\frac{1}{2}, \frac{1}{6}, \frac{1}{6}, \frac{1}{6})$	2.15	-3.14	-2.53
34	$(\frac{1}{2}, \frac{1}{6}, \frac{1}{6}, \frac{1}{6})$	143	$(\frac{1}{6}, \frac{1}{2}, \frac{1}{6}, \frac{1}{6})$	3.06	-4.02	-3.85
35	$(\frac{1}{2}, \frac{1}{6}, \frac{1}{6}, \frac{1}{6})$	143	$(\frac{1}{6}, \frac{1}{6}, \frac{1}{2}, \frac{1}{6})$	3.63	-4.61	-4.60
36	$(\frac{1}{2}, \frac{1}{6}, \frac{1}{6}, \frac{1}{6})$	143	$(\frac{1}{6}, \frac{1}{6}, \frac{1}{6}, \frac{1}{2})$	4.20	-4.93	-5.76
37	$(\frac{1}{6}, \frac{1}{2}, \frac{1}{6}, \frac{1}{6})$	143	$(\frac{1}{2}, \frac{1}{6}, \frac{1}{6}, \frac{1}{6})$	2.35	-2.36	-2.51
38	$(\frac{1}{6}, \frac{1}{2}, \frac{1}{6}, \frac{1}{6})$	143	$(\frac{1}{6}, \frac{1}{2}, \frac{1}{6}, \frac{1}{6})$	3.08	-3.31	-3.68
39	$(\frac{1}{6}, \frac{1}{2}, \frac{1}{6}, \frac{1}{6})$	143	$(\frac{1}{6}, \frac{1}{6}, \frac{1}{2}, \frac{1}{6})$	3.76	-4.57	-5.03
40	$(\frac{1}{6}, \frac{1}{2}, \frac{1}{6}, \frac{1}{6})$	143	$(\frac{1}{6}, \frac{1}{6}, \frac{1}{6}, \frac{1}{2})$	4.48	-5.42	-5.84
41	$(\frac{1}{6}, \frac{1}{6}, \frac{1}{2}, \frac{1}{6})$	143	$(\frac{1}{2}, \frac{1}{6}, \frac{1}{6}, \frac{1}{6})$	2.44	-1.67	-1.91
42	$(\frac{1}{6}, \frac{1}{6}, \frac{1}{2}, \frac{1}{6})$	143	$(\frac{1}{6}, \frac{1}{2}, \frac{1}{6}, \frac{1}{6})$	3.50	-3.00	-3.58
43	$(\frac{1}{6}, \frac{1}{6}, \frac{1}{2}, \frac{1}{6})$	143	$(\frac{1}{6}, \frac{1}{6}, \frac{1}{2}, \frac{1}{6})$	3.95	-4.04	-4.59
44	$(\frac{1}{6}, \frac{1}{6}, \frac{1}{2}, \frac{1}{6})$	143	$(\frac{1}{6}, \frac{1}{6}, \frac{1}{6}, \frac{1}{2})$	4.63	-4.77	-5.72
45	$(\frac{1}{6}, \frac{1}{6}, \frac{1}{6}, \frac{1}{2})$	143	$(\frac{1}{2}, \frac{1}{6}, \frac{1}{6}, \frac{1}{6})$	3.07	-0.86	-2.15
46	$(\frac{1}{6}, \frac{1}{6}, \frac{1}{6}, \frac{1}{2})$	143	$(\frac{1}{6}, \frac{1}{2}, \frac{1}{6}, \frac{1}{6})$	3.97	-2.29	-3.20
47	$(\frac{1}{6}, \frac{1}{6}, \frac{1}{6}, \frac{1}{2})$	143	$(\frac{1}{6}, \frac{1}{6}, \frac{1}{2}, \frac{1}{6})$	4.44	-3.57	-4.47
48	$(\frac{1}{6}, \frac{1}{6}, \frac{1}{6}, \frac{1}{2})$	143	$(\frac{1}{6}, \frac{1}{6}, \frac{1}{6}, \frac{1}{2})$	5.05	-4.26	-5.20
49	$(\frac{1}{4}, \frac{1}{4}, \frac{1}{4}, \frac{1}{4})$	2.5	$(\frac{1}{4}, \frac{1}{4}, \frac{1}{4}, \frac{1}{4})$	8.79	-11.06	-11.43
50	$(\frac{1}{4}, \frac{1}{4}, \frac{1}{4}, \frac{1}{4})$	143	$(\frac{1}{4}, \frac{1}{4}, \frac{1}{4}, \frac{1}{4})$	3.58	-3.60	-4.22

culier car en pratique les valeurs de ces paramètres ne sont pas faciles à collecter.

En considérant la conception d'un essai clinique multicentrique pour lequel un modèle Poisson-gamma sera utilisé pour la modélisation de la dynamique de recrutement,

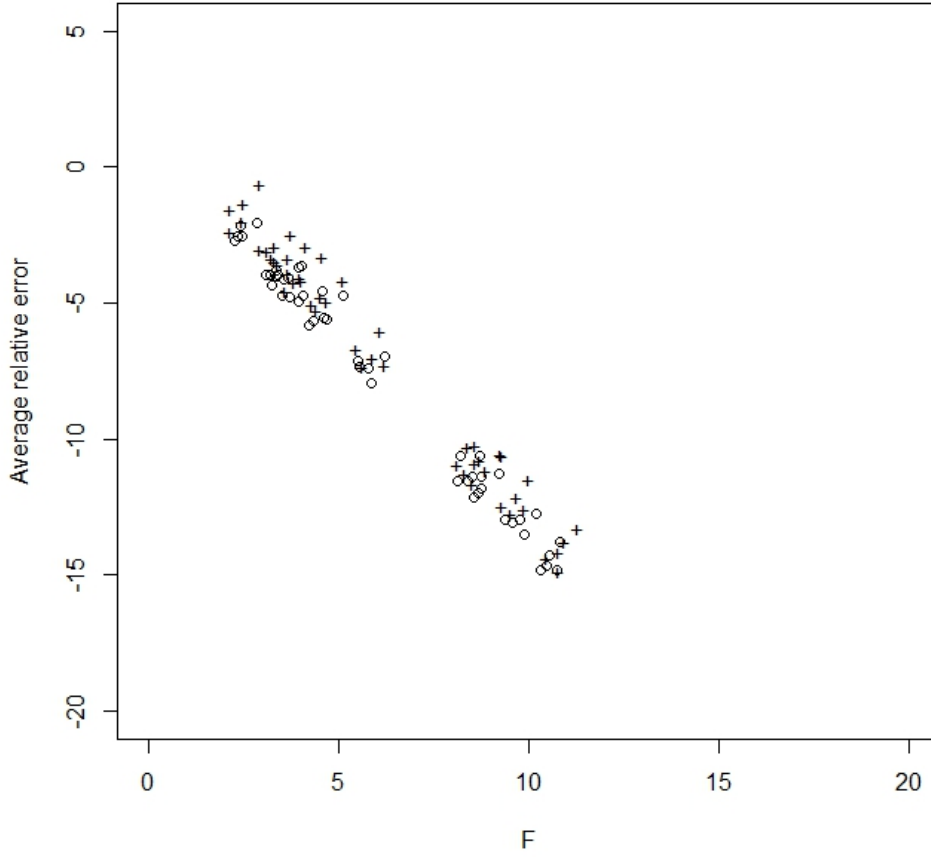


FIGURE 2.8 – Erreur relative moyenne de la prédiction de la durée de recrutement à $t_1 = 78$ semaines, '+' et $t_1 = 104$ semaines, 'o' comme fonction de $\mathbf{F}$.

le nombre de centres C est connu d'avance, une estimation plus ou moins précise des intensités de recrutement est donnée par les investigateurs ainsi que la durée estimée du recrutement. Il est ainsi recommandé lors de cette conception de :

1. Définir le temps intermédiaire t_1 du suivi du recrutement, il est ainsi possible d'estimer les paramètres α, μ du modèle.
2. Vérifier la compatibilité des paramètres intrinsèques à l'essai par une étude telle celle du Scénario 1.
3. Fixer un niveau d'acceptabilité d'erreur relative de prédiction maximum tolérable ε .
4. Effectuer une étude par simulation comme pour les Scénario 3 et 4. Ceci conduit à l'identification de seuils d'acceptabilités de cette perturbations selon ε .

La difficulté vient de la multiplicité des perturbations possibles qui si elles s'ajoutent rendent la définition des seuils d'acceptabilité plus difficile (en effet les perturbations sur l'erreur de prédiction s'ajoutent mais ne sont pas indépendantes quant à leurs effets sur la prédiction).

2.2 Dates d'ouvertures inconnues : Modèle Poisson-Gamma uniforme

Dans le modèle Poisson-Gamma uniforme, les dates d'ouvertures sont supposées uniformément distribuées sur $[0, v_i]$ où v_i est la date de première inclusion du centre i .

Soit $t_1 > 0$. Si le centre i n'a pas recruté sur $[0, t_1]$, on pose $v_i = t_1$. $\tau_i = t_1 - u_i$ est donc uniformément distribué sur $[t_1 - v_i, t_1]$. On pose $S'_i = t_1 - v_i$ et $S_i = t_1$.

Notons que v_i , S_i et S'_i dépendent de t_1 .

Les intensités d'inclusion sont distribuées selon une loi Gamma de paramètres (α, β) .

Proposition 7 *La loi de $k_i = N_i(t_1)$ est donnée par*

$$\mathbb{P}[k_i = n] = \beta^\alpha \frac{\Gamma(\alpha + n)}{\Gamma(\alpha)n!} \int_{S'_i}^{S_i} \frac{t^n}{(\beta + t)^{\alpha+n}} \frac{dt}{S_i - S'_i}.$$

Preuve Sachant l'intensité λ_i et la date d'ouverture du centre u_i , k_i suit une loi de Poisson de paramètre $\lambda_i \tau_i$ où $\tau_i = t_1 - u_i$. τ_i est uniformément distribué sur $[S'_i, S_i]$. On en déduit

$$\begin{aligned} \mathbb{P}[k_i = n] &= \int_0^{+\infty} \beta^\alpha \Gamma(\alpha)^{-1} dx \int_{S'_i}^{S_i} \frac{dt}{S_i - S'_i} e^{-xt} \frac{(xt)^n}{n!} e^{-\beta x} x^{\alpha-1} \\ &= \beta^\alpha \frac{\Gamma(\alpha + n)}{\Gamma(\alpha)n!} \int_{S'_i}^{S_i} \frac{t^n}{(\beta + t)^{\alpha+n}} \frac{dt}{S_i - S'_i}. \end{aligned}$$

2.2.1 Estimation par maximum de vraisemblance

On déduit de la proposition 7 le corollaire :

Corollaire 1 *L'estimateur du maximum de vraisemblance $(\hat{\alpha}_C, \hat{\mu}_C)$ est obtenu en maximisant :*

$$M_C^{\mathcal{U}\Gamma}(\alpha, \mu) = \frac{1}{C} \sum_{i=1}^C [\alpha \ln(\alpha/\mu) - \ln \Gamma(\alpha) + \ln \Gamma(\alpha + k_i) + \ln(B(\alpha, \alpha/\mu, k_i, S'_i, S_i))],$$

où $B(a, b, k, S', S) = \int_{S'+b}^{S+b} t^{-k-a} (t-b)^k dt$.

Remarque Un calcul direct montre que si $a > 1$,

$$B(a, b, k, S', S) = b^{1-a} \left[\beta_{\text{inc}} \left(\frac{b}{S'+b}; a-1, k+1 \right) - \beta_{\text{inc}} \left(\frac{b}{S+b}; a-1, k+1 \right) \right],$$

où β_{inc} est la fonction bêta incomplète : pour tous $x \in [0, 1]$, $a_1 > 0$, $a_2 > 0$,

$$\beta_{\text{inc}}(x; a_1, a_2) = \int_0^x t^{a_1-1} (1-t)^{a_2-1} dt. \quad (2.3)$$

Si $a \neq 1$ et $k > 0$,

$$\begin{aligned} B(a, b, k, S', S) &= \frac{1}{a-1} \left[\frac{S'^k}{(S'+b)^{a+k-1}} - \frac{S^k}{(S+b)^{a+k-1}} \right. \\ &\quad \left. + kb^{1-a} \left(\beta_{\text{inc}} \left(\frac{b}{S'+b}; a, k \right) - \beta_{\text{inc}} \left(\frac{b}{S+b}; a, k \right) \right) \right]. \end{aligned}$$

2.2.2 Propriétés de l'estimateur

Posons $J_i(\theta, k_i, v_i) = -\mathbb{E}_\theta[\nabla^2 f(\theta, k_i, S_i, S'_i)]$ où

$$f(\theta, k_i, S_i, S'_i) = \alpha \ln(\alpha/\mu) - \ln \Gamma(\alpha) + \ln \Gamma(\alpha + k_i) + \ln(B(\alpha, \alpha/\mu, k_i, S'_i, S_i)).$$

Si la suite $(S_i, S'_i)_{i \geq 1}$ vérifie $S_i - S'_i \geq \xi$ pour un certain $\xi > 0$, et si $\frac{1}{C} \sum_{i=1}^C J_i$ converge vers une matrice définie positive I_0 , alors $\hat{\theta}_C$ suit asymptotiquement une loi normale de moyenne θ et de matrice de covariance $\frac{1}{C} I_0^{-1}$.

Ceci nous permet de calculer une région de confiance pour l'estimateur $(\hat{\alpha}_C, \hat{\mu}_C)$: la région de confiance à 95% est une ellipse centrée en $(\hat{\alpha}_C, \hat{\mu}_C)$, définie par

$$\{X \in \mathbb{R}^2 : (X - \theta_0)^T I_0 (X - \theta_0) \leq p/C\},$$

où $p \simeq 5.99$ est le quantile d'ordre 95% d'une loi du χ^2 à deux degrés de liberté.

2.2.3 Validation du modèle

Comme dans le paragraphe 1.5, nous traçons la courbe d'occupation des données, ainsi que la moyenne attendue dans le modèle $\mathcal{UPoisson}-\Gamma$. Cette courbe se fait sur données réelles fournies par l'Inserm, avec l'instant $T_f = 2.32$ années est la durée de recrutement totale. Soit ν_j le nombre de centres ayant recruté exactement j patients, alors

$$\mathbb{E}[\nu_j] = \sum_{i=1}^C \mathbb{P}[k_i = j] = (\alpha/\mu)^\alpha \Gamma(\alpha)^{-1} \Gamma(j + \alpha) \sum_{i=1}^C \frac{1}{S_i - S'_i} B(\alpha, \alpha/\mu, j, S'_i, S_i).$$

La courbe d'occupation est donnée dans la figure 2.9. Le modèle est en bonne adéquation avec les données.

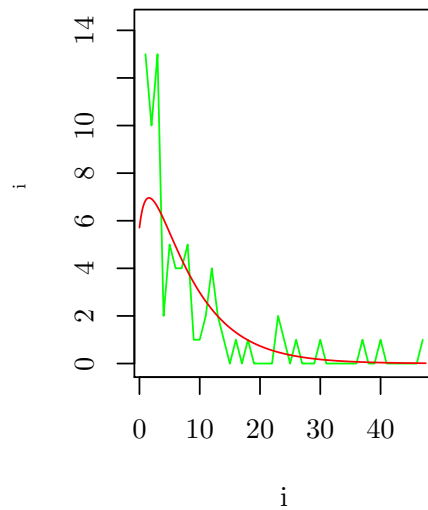


FIGURE 2.9 – Courbe d'occupation à $T_f = 2.31$ années.

2.3 Prise en compte des sorties d'étude

La première situation que nous avons étudiée dans [11] est la prise en compte des "screening failure". Cela consiste à différencier les patients inclus des patients randomisés. En effet, il est normal en recherche clinique d'observer une différence entre ces effectifs due aux critères d'exclusion vérifiés durant la période de screening. Cette différence est problématique car le calcul du NSN (nombre de sujets nécessaires) fournit la valeur du nombre de patients randomisés nécessaires et les données observées servant jusque là au calibrage du modèle sont celles des patients inclus. Usuellement le problème est pris en compte en corrigeant le NSN d'un coefficient - souvent 20 % - pour tenir compte de cette évaporation. Cette méthode est cependant grossière et peut amener (dans un soucis d'atteindre le nombre NSN de patients randomisés) à une phase de recrutement beaucoup plus importante que nécessaire dans le sens où trop de patients finissent par être recrutés (en rappelant que les coûts sont loin d'être négligeables). Ainsi l'apport d'un modèle plus pertinent permet d'estimer au mieux ce coefficient (pouvant dans certaines études être plus grand que 20%).

L'idée est de coupler à la modélisation de la dynamique de recrutement déjà présentée en Section 1.2 une dynamique de sortie d'étude basée sur les mêmes principes de modélisation.

Pour ce faire nous définissons un vecteur de probabilité $(p_1, \dots, p_C)$ représentant les probabilités qu'un patient sorte de l'étude dès l'inclusion (au début de la phase de screening) pour chaque centre.

Le modèle est enrichi en considérant des screening failures à un instant donné durant la période de screening, dont nous supposons que la durée est la même pour tous les patients et de durée strictement positive égale à R .

2.3.1 Introduction

On se place dans les mêmes conditions présentées Section 1.2.

Considérons une étude multicentrique avec C centres. Rappelons que u_i dénote la date d'ouvertures du centre i . Le recrutement des patients est modélisé par des processus doublement stochastiques de Poisson $\{\mathcal{N}_t^i, t \geq 0; i = 1, \dots, C\}$ et le processus global par $\mathcal{N}_t = \sum_{i=1}^C \mathcal{N}_t^i$. Nous notons $\{t_j, j \geq 1\}$ la suite croissante des temps de saut de $\mathcal{N}_t$ (resp. $\{t_{i,j}, j \geq 1\}$ la suite de $\mathcal{N}_t^i$).

Considérons maintenant la modélisation du phénomène de sortie d'étude. Introduisons la famille de variables aléatoires i.i.d $\{r_i; 1 \leq i \leq C\}$ à valeur dans $[0, 1]$, et la famille paramétrique de variables aléatoires positives $\{Z_{i,j}(\theta_i), j \geq 1; 1 \leq i \leq C\}$ où pour chaque i et pour θ_i fixé, les variables $\{Z_{i,j}, j \geq 1\}$ ont la même distribution et les valeurs $\{\theta_i, 1 \leq i \leq C\}$ sont des variables aléatoires i.i.d de distribution quelconque pour le moment.

Par la suite $1 - r_i$ représente la probabilité de sortir de l'étude dès l'inclusion, et $Z_{i,j}(\theta_i)$ le temps de sortie de l'étude après l'inclusion.

Le fait de considérer des variables aléatoires pour r_i et θ_i permet de considérer la variabilité des deux phénomènes de sortie d'étude étudiés ici entre chaque centre.

Le patient j arrivant au centre i au temps $s = t_{i,j}$ peut sortir de l'étude dès l'inclusion avec une probabilité $1 - r_i$, sinon il peut sortir après un temps $s + Z_{i,j}(\theta_i)$ après

l'inclusion pendant la période de screening si $Z_{i,j}(\theta_i) \leq R$, où nous rappelons que R est le temps de screening pour tous les patients. Si aucun de ces événements ne se produit, le patient est alors randomisé avec succès au temps $s + R$.

Nous nous plaçons dans le centre i à un certain temps intermédiaire t_1 . Supposons pour des raisons de simplicité que $R < t_1 - u_i$. Supposons de plus que les valeurs (r_i, θ_i) sont données. Ainsi au temps t_1 pour le patient j recruté au temps $s = t_{i,j} < t_1$, il peut se produire 3 événements :

- Le patient est screené et randomisé avec succès au temps $s + R$ avec probabilité :

$$p_i(s, t_1, r_i, \theta_i) = r_i * \mathbb{P}(Z_{i,j}(\theta_i) > R) \chi(s \leq t_1 - R);$$

- Le patient sort de l'étude avec probabilité :

$$q_i(s, t_1, r_i, \theta_i) = 1 - r_i + r_i \mathbb{P}(Z_{i,j}(\theta_i) \leq \min(R, t_1 - s));$$

- Le patient est toujours en phase de screening avec probabilité :

$$g_i(s, t_1, r_i, \theta_i) = r_i \mathbb{P}(Z_{i,j}(\theta_i) > t_1 - s) \chi(t_1 - R \leq st_1).$$

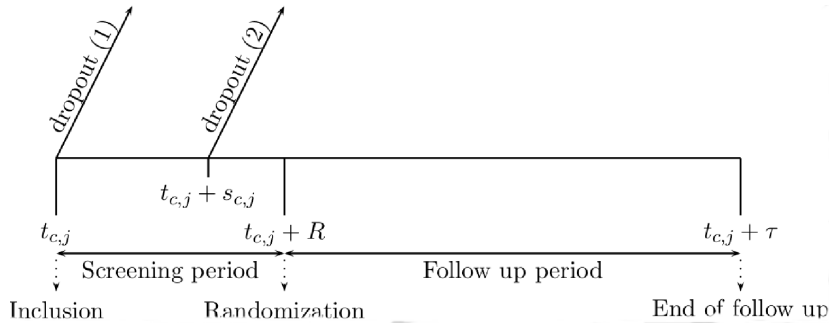


FIGURE 2.10 – Modélisation de la dynamique de recrutement prenant en compte les screening failures.

Définissons la famille indépendante d'indicateurs $\{\chi_{i,j}(r_i), j \geq 1; 1 \leq i \leq C\}$, où pour r_i donné, les variables $\{\chi_{i,j}(r_i), j \geq 1\}$ sont conditionnellement indépendantes et pour tout $1 \leq i \leq C$ et tout $j \geq 1$,

$$\mathbb{P}(\chi_{i,j}(r_i) = 0) = 1 - \mathbb{P}(\chi_{i,j}(r_i) = 1) = r_i$$

Nous pouvons ainsi pour chaque centre i définir trois processus :

- **patients randomisés :**

$$\mathcal{N}_t^{i,R} = \text{Card}\{j : u_i \leq t_{i,j} \leq t_R \text{ et } \chi_{i,j}(r_i) = 0, Z_{i,j}(\theta_i) \geq R\},$$

- **patients perdus :**

$$\mathcal{N}_t^{i,P} = \text{Card}\{j : u_i \leq t_{i,j} \leq t \text{ et } \{\chi_{i,j}(r_i) = 1\} \cup \{Z_{i,j}(\theta_i) \geq \min(R, t - t_{i,j})\}\},$$

— patients en cours de screening :

$$\mathcal{N}_t^{i,S} = \mathcal{N}_t^i - \mathcal{N}_t^{i,R} - \mathcal{N}_t^{i,P}$$

Finalement notons

$$\mathcal{N}_t^X = \sum_{i=1}^C \mathcal{N}_t^{i,X},$$

pour $X := R; P, S$.

2.3.2 Modèles de prise en compte des sorties d'étude

La période de recrutement s'arrête lorsque $\mathcal{N}_t^{i,R} = N$. Plusieurs modèles de sortie d'étude sont étudiés par la suite :

Modèle 1 : Pour tout $1 \leq i \leq C$, $r_i = r$, où r est une constante fixée dans $[0, 1]$ et pour tout $j \geq 1$, $Z_{i,j}(\theta_i) = +\infty$. Cela signifie que seules les sorties d'étude dès l'inclusion sont considérées.

Modèle 2 : Les variables $\{r_i; 1 \leq i \leq C\}$ sont i.i.d dont la distribution est une distribution beta de paramètres (ψ_1, ψ_2) , et pour tout $1 \leq i \leq C$ et $j \geq 1$, $Z_{i,j}(\theta_i) = +\infty$. De nouveau seules les sorties d'étude dès l'inclusion sont considérées mais un ajout de variabilité dans les probabilités de randomisation entre centre est considéré par l'ajout de la distribution beta.

Modèle 3 : Pour tout $1 \leq i \leq C$, $r_i = r$, et les valeurs $\{Z_{i,j}(\cdot), j \geq 1; 1 \leq i \leq C\}$ sont des variables i.i.d de distribution exponentielle de paramètres θ (le même pour tout les centres).

Modèle 4 : Pour tout $1 \leq i \leq C$, $r_i = r$, et les valeurs $\{Z_{i,j}(\cdot), j \geq 1\}$ sachant θ_i sont des variables i.i.d. de distribution exponentielle de paramètre θ_i , où les valeurs $\{\theta_i; 1 \leq i \leq C\}$ sont des variables i.i.d. de distribution gamma de paramètres (α_2, μ_2) .

Modèle 5 : Les variables $\{r_i; 1 \leq i \leq C\}$ sont i.i.d et dont la distribution est une distribution beta de paramètres (ψ_1, ψ_2) . Les valeurs $\{Z_{i,j}(\cdot), j \geq 1\}$ sachant θ_i sont des variables i.i.d. de distribution exponentielle de paramètre θ_i , où les valeurs $\{\theta_i; 1 \leq i \leq C\}$ sont des variables i.i.d. de distribution gamma de paramètres (α_2, μ_2) .

Le modèle 5 est le modèle le plus complexe qui considère les deux types de sortie d'étude en prenant en compte la variabilité de randomisation entre les centres pour chacun des types.

Rappelons que la dynamique de recrutement est modélisée par un processus Poisson-gamma. La démarche est identique à celle développée au paragraphe 1 à savoir les résultats reposent sur la ré-estimation bayésienne empirique des paramètres des familles de loi paramétriques choisies que sont les lois Beta et les lois gamma intervenant dans le

modèle. Il est important de noter que les lois inhérentes à la dynamique de recrutement sont indépendantes de celles de la dynamique de sortie d'étude. Comme au paragraphe 1, la problématique vient du fait que ces paramètres sont inconnus mais estimables au moyen de données issues d'analyse intermédiaire de recrutement et de sortie d'étude dans ce cas précis. On est alors en mesure de construire un modèle prédictif du nombre de patients randomisés (nombre de patients inclus - nombre de drop out des deux types).

Dans les modèles 1 et 2, le véritable temps de sortie d'étude n'est pas considéré. Ainsi au temps intermédiaire t_1 et pour chaque centre i , seules les données du nombre de patients recrutés et randomisés sont nécessaires (i.e $\{N_{t_1}^i, N_{t_1}^{i,R}\}$). En ce qui concerne les modèles 3, 4 et 5, nous supposons que toutes les données sont observées et récoltées : pour chaque patient sont connus le temps de recrutement ainsi que le temps de sortie d'étude (ou de randomisation). Sans la donnée du temps de sortie d'étude aucune distinction n'est possible entre une sortie d'étude dès l'inclusion et une sortie d'étude pendant la période de screening, et la distinction faite entre les modèles serait inutile.

2.3.3 Estimation des paramètres au temps intermédiaire t_1 .

Le recrutement est toujours modélisé par un modèle Poisson-gamma comme présenté au Chapitre 1. Nous supposons toujours les intensités de recrutement des centres comme distribuées selon une loi Gamma de paramètres inconnus (α, β) .

Soit t_1 le temps intermédiaire d'étude, considérons le cas du centre i , et supposons pour des raisons de simplification que $\tau_i = t_1 - u_i \geq R$. Dans cette situation le nombre de patients recrutés au temps t_1 est $n_i = N_{t_1}^i$, est vu comme une variable aléatoire de distribution une loi binomiale négative de paramètres $(\alpha, \tau_i/(\beta + \tau_i))$ [38, p.199] (ou une variable aléatoire de distribution Poisson-gamma de paramètres (α, β, τ_i) [16, p.119]). Rappelons que la distribution de probabilité d'une variable binomiale négative (Poisson-gamma) de paramètres (α, π) s'écrit :

$$\text{NegBin}(k; \alpha, \pi) = \frac{\Gamma(\alpha + k)}{k! \Gamma(\alpha)} \pi^k (1 - \pi)^\alpha.$$

Notons $\mu = \mathbb{E}[\lambda]$ la moyenne des intensités de recrutement et par $\sigma^2 = \mathbb{V}[\lambda]$ leur variance en rappelant que : $\mu = \alpha/\beta$ et $\sigma^2 = \alpha/\beta^2$. Ainsi $\alpha = \mu^2/\sigma^2$ et $\beta = \mu/\sigma^2$. En reparamétrisant en terme de moyenne d'intensité pour le centre i on obtient la probabilité,

$$\mathbb{P}(n_i = k) = \mathbb{E} \left[e^{-\lambda \tau_i} \frac{(\lambda \tau_i)^k}{k!} \right] = \text{NegBin} \left(k; \alpha, \frac{\mu \tau_i}{\alpha + \mu \tau_i} \right). \quad (2.4)$$

Modèle 1.

Soit r_i la probabilité de randomisation du centre i . Notons $\text{Bin}(n, \pi)$ une variable aléatoire binomiale de paramètres (n, π) dont la distribution de probabilité est :

$$\text{Bin}(k; n, \pi) = \binom{n}{k} \pi^k (1 - \pi)^{n-k}, \quad 0 \leq k \leq n.$$

Supposons pour simplifier qu'il n'y a pas de délai dans la phase de screening ($R = 0$). Ainsi pour le centre i le nombre de patients randomisés noté k_i a une distribution

binomiale de paramètres (n_i, r_i) . Supposons maintenant que $r_i \equiv r$ (les probabilités de randomisation sont les mêmes pour tous les centres). Ainsi, ayant toutes les données des C centres, la fonction de log-vraisemblance est facilement calculable.

Sachant $\{n_i, k_i, \tau_i\}$ au temps intermédiaire t_1 , la fonction log-vraisemblance s'écrit sous la forme :

$$\mathcal{L}_1(\alpha, \mu, r) = \sum_{i=1}^C \ln \left[\text{NegBin} \left(n_i; \alpha, \frac{\mu \tau_i}{\alpha + \mu \tau_i} \right) \right] + \sum_{i=1}^C \ln [\text{Bin}(k_i; n_i, r)].$$

On voit que le paramètre r est séparé des paramètres (α, μ) et $\mathcal{L}_1(\alpha, \mu, r)$ peut se réécrire sous la forme : $\mathcal{L}_1(\alpha, \mu, r) = \mathcal{L}_{1,1}(\alpha, \mu) + \mathcal{L}_{1,2}(r)$, avec,

$$\mathcal{L}_{1,1}(\alpha, \mu) = \sum_{i=1}^C \ln \Gamma(n_i + \alpha) - C \ln \Gamma(\alpha) + N_1 (\ln \mu - \ln \alpha) - \sum_{i=1}^C (n_i + \alpha) \ln(1 + \mu \tau_i / \alpha) + C, \quad (2.5)$$

et

$$\mathcal{L}_{1,2}(r) = \sum_{i=1}^C [k_i \ln r + (n_i - k_i) \ln(1 - r)] + C_1.$$

où C_1 est une constante indépendante des paramètres considérés, et $N_1 = \sum_{i=1}^C n_i$ est le nombre total de patients recrutés jusqu'au temps t_1 . En dérivant par rapport à r , il est facile de voir que l'estimateur du maximum de vraisemblance de la fonction $\mathcal{L}_{1,2}$ est donné par l'équation suivante :

$$\hat{r} = \left(\sum_{i=1}^M n_i \right)^{-1} \sum_{i=1}^M k_i. \quad (2.6)$$

Remarque L'indépendance du processus de sortie d'étude et du processus de recrutement implique que, pour tout les modèles considérés, les paramètres (α, μ) peuvent être estimés par la fonction de log-vraisemblance $\mathcal{L}_{1,1}$ donnée (2.5) et une procédure d'optimisation à 2 dimensions. Nous notons $(\hat{\alpha}, \hat{\mu})$ ces estimations

Remarque Notons de plus que l'estimateur de la variance est donné par $\hat{\sigma}^2 = \hat{\mu}^2 / \hat{\alpha}$.

Modèle 2.

Supposons que les r_i peuvent varier entre les centres. Cette variation peut être décrite en considérant le vecteur $(r_i, 1 \leq i \leq C)$ comme distribué selon une loi beta de paramètres inconnus (ψ_1, ψ_2) .

Notons $\text{Beta}(\psi_1, \psi_2)$ une variable aléatoire de distribution beta dont la fonction de densité s'écrit :

$$p_\beta(x; \psi_1, \psi_2) = x^{\psi_1-1} (1-x)^{\psi_2-1} / \mathcal{B}(\psi_1, \psi_2), \quad x \in]0, 1[,$$

où $\mathcal{B}(\psi_1, \psi_2) = \int_0^1 x^{\psi_1-1} (1-x)^{\psi_2-1} dx$ est la fonction Beta.

Notons que si $r = \text{Beta}(\psi_1, \psi_2)$, alors une variable binomiale doublement stochastique $\text{Bin}(n, r)$ a une distribution beta-binomiale dont la distribution de probabilité est donnée par :

$$\mathbb{P}(\text{Bin}(n; \text{Beta}(\psi_1, \psi_2)) = k) = \binom{n}{k} \frac{\mathcal{B}(k + \psi_1, n - k + \psi_2)}{\mathcal{B}(\psi_1, \psi_2)}.$$

Sachant $\{n_i, k_i, \tau_i\}$ au temps intermédiaire t_1 , la fonction log-vraisemblance peut s'écrire sous la forme :

$$\mathcal{L}_2(\alpha, \mu, \psi_1, \psi_2) = \mathcal{L}_{2,1}(\alpha, \mu) + \mathcal{L}_{2,2}(\psi_1, \psi_2),$$

où $\mathcal{L}_{2,1} = \mathcal{L}_{1,1}$ est donnée par (2.5), et

$$\mathcal{L}_{2,2}(\psi_1, \psi_2) = \sum_{i=1}^M \ln \mathcal{B}(k_i + \psi_1, n_i - k_i + \psi_2) - M \ln \mathcal{B}(\psi_1, \psi_2) + C_2. \quad (2.7)$$

Les paramètres (ψ_1, ψ_2) peuvent être estimés en utilisant la fonction log-vraisemblance $\mathcal{L}_{2,2}$ donnée par (2.7) et une procédure d'optimisation à 2 dimensions. Notons les estimations de la fonction $\mathcal{L}_{2,2}$ par $(\hat{\psi}_1, \hat{\psi}_2)$. Considérons maintenant une approche bayésienne d'ajustement des paramètres pour chaque centre i sachant n_i et k_i de la même manière que dans [8, 6] et dans le Chapitre 1. Rappelons que la distribution *a posteriori* des intensités de recrutement est une distribution gamma de paramètres $(\hat{\alpha} + n_i, \hat{\beta} + \tau_i)$. De la même manière, si r_i a une distribution *a priori* beta de paramètres (ψ_1, ψ_2) , alors sachant $\{n_i, k_i\}$, la distribution *a posteriori* des r_i est aussi une distribution beta de paramètres $(\psi_1 + k_i, \psi_2 + n_i - k_i)$.

Ainsi, ayant toutes les données nécessaires, il est possible de représenter les estimateurs *a posteriori* des intensités et des probabilités de randomisation pour chaque centre sous la forme :

$$\hat{\lambda}_i \sim \Gamma(\hat{\alpha} + n_i, \hat{\beta} + \tau_i) \quad \text{et} \quad \hat{r}_i \sim \text{Beta}(\hat{\psi}_1 + k_i, \hat{\psi}_2 + n_i - k_i), \quad (2.8)$$

où $\hat{\lambda}_i$ et $\hat{r}_i, i = 1, \dots, M$ sont indépendants.

Modèles 3, 4 et 5.

Lors du calcul de la fonction vraisemblance, il est important de noter que les variables $\{\lambda_i, r_i, \theta_i ; 1 \leq i \leq C\}$ sont indépendantes. Nous supposons toujours que les intensités $\{\lambda_i ; 1 \leq i \leq C\}$ sont distribuées selon une loi gamma de paramètres (α, β) , ainsi pour tout $1 \leq i \leq C$ la variable N_t^i a une distribution binomiale négative $\text{NegBin}\left(k; \alpha, \frac{\mu\tau_i}{\alpha + \mu\tau_i}\right)$ (voir (2.4)). Le type de distribution considéré pour r_i et θ_i est spécifié par le choix du modèle considéré.

Nous supposons dans ces modèles que nous disposons de plus de données, ainsi il est possible d'estimer les paramètres des temps de sortie d'étude $\{Z_{i,j}, j \geq 1 ; 1 \leq i \leq C\}$. De plus, le calcul des distribution *a posteriori* des paramètres n'est pas simple si l'on utilise les données comme pour les modèles 1 et 2. Au temps t_1 , on observe les temps d'arrivée des patients $\{t_{i,j} \leq t_1 ; j \geq 1, 1 \leq i \leq C\}$ et le dernier temps $\{s_{i,j}, j \geq 1 ; 1 \leq i \leq C\}$ auquel il était en phase de screening (i.e $t_{i,j} \leq s_{i,j} \leq t_1$). Ceci implique que l'on observe également $\{\min(Z_{i,j}, R, t_1 - t_{i,j}) \vee 0, j \geq 1 ; 1 \leq i \leq C\}$, où comme précédemment $a \vee b = \max(a, b)$.

Notons m_i le nombre de patients perdus lors de la phase de screening du centre i , $T_i = \sum_j (s_{i,j} - t_{i,j})$ - la somme des durées de screening du centre i , et $\tilde{k}_i$ le nombre de patients qui ne sont pas perdus dès l'inclusion.

Sachant $\{\theta_i, r_i ; 1 \leq i \leq C\}$, on peut écrire une formule générale pour la vraisemblance,

$$\mathcal{L}[(t_{i,j}); (s_{i,j})] = \exp[\mathcal{L}_{1,1}(\alpha, \mu)] \times \prod_{i=1}^C [\mathcal{L}_2(\theta_i, r_i; (t_{i,j}), (s_{i,j}))] \quad (2.9)$$

où $\mathcal{L}_{1,1}$ est donnée par (2.5), et

$$\begin{aligned} \mathcal{L}_2(\theta_i, r_i; (t_{i,j}), (s_{i,j})) &= \prod_{\mathcal{D}_1} (1 - r_i) \prod_{\mathcal{D}_2} r_i \exp(-\theta_i(s_{i,j} - t_{i,j})) \\ &\quad \times \prod_{\mathcal{D}_3} r_i \theta_i \exp(-\theta_i(s_{i,j} - t_{i,j})), \end{aligned}$$

où $\mathcal{D}_1 = \bigcup_{i=1}^C \mathcal{D}_1^i$, $\mathcal{D}_2 = \bigcup_{i=1}^C \mathcal{D}_2^i$, $\mathcal{D}_3 = \bigcup_{i=1}^C \mathcal{D}_3^i$ et pour tout $1 \leq i \leq C$,

$$\begin{aligned} \mathcal{D}_1^i &= \{j \geq 1, \text{ s.t. } s_{i,j} = t_{i,j}\}, \\ \mathcal{D}_2^i &= \{j \geq 1, \text{ s.t. } s_{i,j} = (t_{i,j} + R) \wedge t_1\}, \\ \mathcal{D}_3^i &= \{j \geq 1, \text{ s.t. } t_{i,j} < s_{i,j} < (t_{i,j} + R) \wedge t_1\}. \end{aligned}$$

Dans (2.9), l'espérance est calculée selon θ_i et r_i qui varient selon leurs distributions respectives définies par le modèle considéré.

Notons que $m_i = \text{Card}(\mathcal{D}_3^i)$ et $\tilde{k}_i = \text{Card}(\mathcal{D}_2^i) + \text{Card}(\mathcal{D}_3^i)$. Ainsi $\mathcal{L}_2$ peut se réécrire comme,

$$\mathcal{L}_2(\theta_i, r_i; (t_{i,j}), (s_{i,j})) = (1 - r_i)^{n_i - \tilde{k}_i} r_i^{\tilde{k}_i} \theta_i^{m_i} \exp(-\theta_i T_i). \quad (2.10)$$

Modèle 3.

Les estimateurs du maximum de vraisemblance de r et θ sont donnés par :

$$\hat{r} = \left(\sum_{i=1}^M n_i \right)^{-1} \sum_{i=1}^M \tilde{k}_i \quad \text{et} \quad \hat{\theta} = \left(\sum_{i=1}^M T_i \right)^{-1} \sum_{i=1}^M m_i, \quad (2.11)$$

En effet, sachant $\{t_{i,j}, s_{i,j}, \tau_i, j \geq 1; 1 \leq i \leq C\}$ au temps intermédiaire t_1 , et en utilisant (2.9) et (2.10), la fonction log-vraisemblance peut se réécrire comme :

$$\mathcal{L}_3(\alpha, \mu, r, \theta) = \mathcal{L}_{3,1}(\alpha, \mu) + \sum_{i=1}^C \left[(n_i - \tilde{k}_i) \ln(1 - r) + \tilde{k}_i \ln r + m_i \ln \theta - \theta T_i \right],$$

où $\mathcal{L}_{3,1} = \mathcal{L}_{1,1}$ est donné par (2.5). En prenant les dérivées selon r et θ on obtient les estimateurs du maximum de vraisemblance de (2.11).

Modèle 4.

Sachant $\{t_{i,j}, s_{i,j}, \tau_i, j \geq 1; 1 \leq i \leq C\}$ au temps intermédiaire t_1 , et en utilisant (2.9) et (2.10), la fonction de log-vraisemblance peut se réécrire sous la forme :

$$\mathcal{L}_4(\alpha, \mu, r, \alpha_2, \beta_2) = \mathcal{L}_{4,1}(\alpha, \mu) + \sum_{i=1}^C \left[(n_i - \tilde{k}_i) \ln(1 - r) + \tilde{k}_i \ln r \right] + \mathcal{L}_{4,2}(\alpha_2, \beta_2)$$

où

$$\mathcal{L}_{4,2}(\alpha_2, \beta_2) = \sum_{i=1}^C \left[\ln \Gamma(m_i + \alpha_2) - \ln \Gamma(\alpha_2) + \alpha_2 \ln \beta_2 - (m_i + \alpha_2) \ln(\beta_2 + T_i) \right] \quad (2.12)$$

et $\mathcal{L}_{4,1} = \mathcal{L}_{1,1}$ est donnée par (2.5). $(\hat{\alpha}, \hat{\mu})$ et $(\hat{\alpha}_2, \hat{\beta}_2)$ sont calculés numériquement et par dérivation $\hat{r}$ est donné par (2.11).

Les paramètres (α_2, β_2) peuvent être estimés en utilisant la fonction log-vraisemblance $\mathcal{L}_{4,2}$ donnée par (2.12).

En considérant une approche bayésienne de reparamétrisation pour chaque centre sachant les données $\{n_i, \tau_i, (t_{i,j}), (s_{i,j}), j \geq 1 ; 1 \leq i \leq C\}$, on a θ_i qui a une distribution *a priori* gamma de paramètres (α_2, β_2) . Sachant $\{m_i, T_i ; 1 \leq i \leq C\}$ (la connaissance de $\{t_{i,j}\}$ et $\{s_{i,j}\}$ n'est pas nécessaire ici) et en utilisant les formules bayésiennes, on obtient la distribution *a posteriori* de θ_i qui est une distribution gamma de paramètres $(\alpha_2 + m_i, \beta_2 + T_i)$. Ainsi,

$$\begin{aligned}\hat{\lambda}_i &\sim \Gamma(\hat{\alpha} + n_i, \hat{\beta} + \tau_i), \\ \hat{r} &= \left(\sum_{i=1}^C n_i \right)^{-1} \sum_{i=1}^C \tilde{k}_i, \\ \hat{\theta}_i &\sim \Gamma(\hat{\alpha}_2 + m_i, \hat{\beta}_2 + T_i), \quad i = 1, \dots, C.\end{aligned}\tag{2.13}$$

Modèle 5.

Sachant $\{t_{i,j}, s_{i,j}, \tau_i, j \geq 1 ; 1 \leq i \leq C\}$ au temps intermédiaire t_1 , et en utilisant (2.9) et (2.10), la fonction log-vraisemblance peut se réécrire comme :

$$\mathcal{L}_5(\alpha, \mu, \psi_1, \psi_2, \alpha_2, \beta_2) = \mathcal{L}_{5,1}(\alpha, \mu) + \mathcal{L}_{5,2}(\alpha_2, \beta_2) + \mathcal{L}_{5,3}(\psi_1, \psi_2)$$

où

$$\mathcal{L}_{5,3}(\psi_1, \psi_2) = \sum_{i=1}^M \left[\ln \mathcal{B}(\tilde{k}_i + \psi_1, n_i - \tilde{k}_i + \psi_2) - \ln \mathcal{B}(\psi_1, \psi_2) \right], \tag{2.14}$$

$\mathcal{L}_{5,1} = \mathcal{L}_{1,1}$ est donnée par (2.5) et $\mathcal{L}_{5,2} = \mathcal{L}_{4,2}$ est donnée par (2.12).

De manière similaire à (2.13), θ_i a une distribution *a priori* gamma de paramètres (α_2, β_2) , ainsi sachant $\{m_i, T_i ; 1 \leq i \leq C\}$, la distribution *a posteriori* de θ_i est une distribution gamma de paramètres $(\alpha_2 + m_i, \beta_2 + T_i)$. De plus, r_i a une distribution *a priori* beta de paramètres (ψ_1, ψ_2) , ainsi sachant $\{n_i, \tilde{k}_i ; 1 \leq i \leq C\}$, la distribution *a posteriori* de r_i est une distribution beta de paramètres $(\psi_1 + \tilde{k}_i, \psi_2 + n_i - \tilde{k}_i)$. Pour conclure, les distributions *a posteriori* des intensités d'inclusion, des probabilités de sortie d'étude dès l'inclusion et les intensités de sortie d'étude pendant la phase de screening sont données par :

$$\hat{\lambda}_i \sim \Gamma(\hat{\alpha} + n_i, \hat{\beta} + \tau_i), \tag{2.15}$$

$$\hat{r}_i = \text{Beta}(\hat{\psi}_1 + \tilde{k}_i, \hat{\psi}_2 + n_i - \tilde{k}_i), \tag{2.16}$$

$$\hat{\theta}_i \sim \Gamma(\hat{\alpha}_2 + m_i, \hat{\beta}_2 + T_i), \quad i = 1, \dots, C. \tag{2.17}$$

2.3.4 Prédiction.

Prédiction du nombre de patients randomisés.

Pour l'ensemble des modèles nous notons $K_2 = \sum_{i=1}^C k_i$.

Modèle 1. Sachant $\{(n_i, \tau_i) ; 1 \leq i \leq C\}$ au temps intermédiaire t_1 , la prédiction du nombre de patients recrutés $\hat{n}_i(t)$ par le centre i est un processus Poisson-gamma d'intensités *a posteriori* $\hat{\lambda}_i \sim \Gamma(\hat{\alpha} + n_i, \hat{\beta} + \tau_i)$ (voir (2.8) et [6]). Considérons maintenant

la prédiction du nombre de patients qui vont être randomisés. Notons $(\hat{\mu}, \hat{\beta}, \hat{r})$ les estimateurs des paramètres (μ, β, r) . Rappelons que R est la durée de la période de screening. Posons ν_i représentant le nombre de patients qui sont entrés dans la phase de screening pour le centre i dans l'intervalle temporel $[t_1 - R, t_1]$ et k_i le nombre total de patients randomisés jusqu'au temps t_1 .

Ainsi, sachant ν_i , le nombre de patients qui vont être randomisés dans l'intervalle $[t_1, t_1 + R]$ est une variable aléatoire binomiale $\text{Bin}(\nu_i, r)$ et les temps auxquels ces patients sont randomisés sont distribués uniformément dans l'intervalle $[t_1, t_1 + R]$. La prédiction du nombre de patients randomisés dans l'intervalle $[t_1, t_1 + R]$ est une variable aléatoire binomiale $\text{Bin}(\nu_i, \hat{r})$, où $\hat{r}$ est défini par (2.6). Le nombre de patients randomisés après le temps $t_1 + R$ peut être considéré comme un amincissement du processus de Poisson N_i^t avec probabilité r . Posons Π_a représentant un processus de Poisson d'intensité a . Ainsi, pour tout $t > t_1 + R$, le processus prédictif du nombre de patients randomisés par le centre i $\{\hat{k}_i(t), t \geq t_1 + R\}$, peut être vu comme un processus de Poisson d'intensité $\hat{r}\hat{\lambda}_i$. Donc,

$$\hat{k}_i(t) = k_i + \text{Bin}(\nu_i, \hat{r}) + \Pi_{\hat{r}\hat{\lambda}_i}(t - t_1 - R). \quad (2.18)$$

Rappelons que pour une intensité aléatoire λ , $\mathbb{E}[\Pi_\lambda(t)] = t\mathbb{E}[\lambda]$, $\mathbb{V}[\Pi_\lambda(t)] = t\mathbb{E}[\lambda] + t^2\mathbb{E}[\lambda]$. Ainsi, selon les données disponibles,

$$\mathbb{E}[\hat{k}_i(t) \mid \text{data}] = k_i + \nu_i\hat{r} + \hat{r}(t - t_1 - R)\mathbb{E}[\hat{\lambda}_i],$$

où $\hat{r}$ est donnée par (2.6), et

$$\mathbb{E}[\hat{\lambda}_i] = \frac{\alpha + n_i}{\beta + \tau_i}, \quad \mathbb{V}[\hat{\lambda}_i] = \frac{\alpha + n_i}{(\beta + \tau_i)^2}. \quad (2.19)$$

Sachant que si l'on connaît $\{n_i, k_i\}$, les estimations *a posteriori* $\hat{\lambda}_i$ sont indépendants de $\hat{r}$, alors la variance est

$$\mathbb{V}[\hat{k}_i(t) \mid \text{data}] = \nu_i\hat{r}(1 - \hat{r}) + \hat{r}(t - t_1 - R)\mathbb{E}[\hat{\lambda}_i] + \hat{r}^2(t - t_1 - R)^2\mathbb{V}[\hat{\lambda}_i]. \quad (2.20)$$

Finalement,

$$\begin{aligned} \mathbb{E}[\hat{k}(t) \mid \text{data}] &= K_2 + \hat{r} \sum_{i=1}^C \nu_i + \hat{r}(t - t_1 - R) \sum_{i=1}^C \frac{\alpha + n_i}{\beta + \tau_i}, \\ \mathbb{V}[\hat{k}(t) \mid \text{data}] &= \hat{r}(1 - \hat{r}) \sum_{i=1}^C \nu_i + \hat{r}(t - t_1 - R) \sum_{i=1}^C \frac{\alpha + n_i}{\beta + \tau_i} \\ &\quad + \hat{r}^2(t - t_1 - R)^2 \sum_{i=1}^C \frac{\alpha + n_i}{(\beta + \tau_i)^2}. \end{aligned}$$

Modèle 2. Pour ce modèle, pour $t > t_1 + R$ il est possible de considérer pour le processus prédictif $\hat{k}_i(t)$ du centre i une formule similaire à (2.18), où les estimations *a posteriori* $\hat{\lambda}_i$ et $\hat{r}_i$ sont donnés par (2.8).

Modèles 3, 4, 5. Pour les modèles 3, 4, et 5, étant donné que l'on dispose de plus d'informations au temps t_1 , le cas d'un patient où le résultat de la phase de screening est inconnue correspond au cas $t_1 - R < t_{i,j} \leq t_1$ et $Y_{i,j} > t_1 - t_{i,j}$. Pour $1 \leq i \leq C$, posons Ω_i l'ensemble des indices,

$$\Omega_i = \{j \in \mathbb{N} : t_1 - R < t_{i,j} \leq t_1 \text{ et } Y_{i,j} > t_1 - t_{i,j}\} \quad (2.21)$$

et $\nu_i = \text{Card}(\Omega_i)$. Conditionnellement à θ_i , la probabilité qu'à le patient d'être randomisé dans l'intervalle $[t_1, t_1 + R]$ est

$$\mathbb{P}Y_{i,j} \geq R | Y_{i,j} > t_1 - t_{i,j}; \theta_i = e^{-\theta_i(t_{i,j}+R-t_1)}$$

Sachant les données jusqu'à t_1 et θ_i , le nombre de patients randomisés entre t_1 et $t_1 + R$ pour le centre i est la somme de ν_i variables aléatoires de Bernoulli de probabilité $e^{-\theta_i(R-t_1+t_{i,j})}$ notée $\text{Ber}(e^{-\theta_i(t_{i,j}+R-t_1)})$. Ainsi, le processus prédictif $\{\hat{k}_i(t), t \geq t_1 + R\}$ pour $t > t_1 + R$ peut se réécrire comme

$$\hat{k}_i(t) = k_i + \Pi_{\hat{p}_i \hat{\lambda}_i}(t - t_1 - R) + \sum_{j \in \Omega_i} \text{Ber}(e^{-\theta_i(t_{i,j}+R-t_1)}). \quad (2.22)$$

La probabilité de non sortie d'étude est $\hat{p}_i = \hat{r}_i \exp(-\hat{\theta}_i R)$.

La moyenne et variance de λ_i est calculée dans (2.19), et

$$\mathbb{E}[\hat{r}_i | data] = \frac{\psi_1 + k_i}{\psi_1 + \psi_2 + n_i}, \quad \mathbb{V}[\hat{r}_i | data] = \frac{(\psi_1 + k_i)(\psi_2 + n_i - k_i)}{(\psi_1 + \psi_2 + n_i)^2(1 + \psi_1 + \psi_2 + n_i)},$$

où à la place de (ψ_1, ψ_2) est utilisé $(\hat{\psi}_1, \hat{\psi}_2)$. Sachant $\{n_i, k_i; 1 \leq i \leq C\}$, les estimations *a posteriori* $\hat{\lambda}_i$ et $\hat{r}_i$ sont indépendants. Rappelons que pour une probabilité aléatoire r , $\mathbb{E}[\text{Bin}(n, r)] = n\mathbb{E}[r]$, et

$$\mathbb{V}[\text{Bin}(n, r)] = \mathbb{E}[\mathbb{V}[\text{Bin}(n, r) | r]] + \mathbb{V}[\mathbb{E}[\text{Bin}(n, r) | r]] = n\mathbb{E}[r(1-r)] + n^2\mathbb{V}[r].$$

Ainsi,

$$\mathbb{E}[\hat{k}_i(t) | data] = k_i + \nu_i \mathbb{E}[\hat{r}_i] + (t - t_1 - R) \mathbb{E}[\hat{r}_i] \mathbb{E}[\hat{\lambda}_i],$$

et

$$\mathbb{V}[\hat{k}_i(t) | data] = \nu_i \mathbb{E}[\hat{r}_i(1 - \hat{r}_i)] + \nu_i^2 \mathbb{V}[\hat{r}_i] + (t - t_1 - R) \mathbb{E}[\hat{r}_i] \mathbb{E}[\hat{\lambda}_i] + (t - t_1 - R)^2 \mathbb{V}[\hat{r}_i \hat{\lambda}_i],$$

où l'on peut utiliser la formule,

$$\mathbb{V}[\hat{r}_i \hat{\lambda}_i] = \mathbb{E}[\hat{\lambda}_i^2] \mathbb{V}[\hat{r}_i] + (\mathbb{E}[\hat{r}_i])^2 \mathbb{V}[\hat{\lambda}_i].$$

Finalement en utilisant les relations (2.19), (2.20) on peut calculer la moyenne et la variance du processus global $\hat{k}(t)$.

Modèle 3. Ici $\hat{r}_i \equiv \hat{r}$ et $\hat{\theta}_i \equiv \hat{\theta}$, où $(\hat{r}, \hat{\theta})$ sont donnés par (2.11). Ainsi en utilisant (2.22) on obtient

$$\hat{k}_i(t) = k_i + \Pi_{\hat{p} \hat{\lambda}_i}(t - t_1 - R) + \sum_{j \in \Omega_i} \text{Ber}(e^{-\hat{\theta}(t_{i,j}+R-t_1)}),$$

où $\hat{p} = \hat{r}e^{-\hat{\theta}R}$.

Nous avons,

$$\begin{aligned} \mathbb{E}[\hat{k}_i(t) | data] &= k_i + (t - t_1 - R) \hat{r} e^{-\hat{\theta}R} \mathbb{E}[\hat{\lambda}_i] + \sum_{j \in \Omega_i} e^{-\hat{\theta}(t_{i,j}+R-t_1)}, \\ \mathbb{V}[\hat{k}_i(t) | data] &= (t - t_1 - R)^2 \hat{r}^2 e^{-2\hat{\theta}R} \mathbb{V}[\hat{\lambda}_i] + (t - t_1 - R) \hat{r} e^{-\hat{\theta}R} \mathbb{E}[\hat{\lambda}_i] \\ &\quad + \sum_{j \in \Omega_i} e^{-\hat{\theta}(t_{i,j}+R-t_1)} (1 - e^{-\hat{\theta}(t_{i,j}+R-t_1)}), \end{aligned}$$

où $\mathbb{E}[\hat{\lambda}_i]$ et $\mathbb{V}[\hat{\lambda}_i]$ sont donnés par (2.19).

Modèle 4. Le processus prédictif du centre i est

$$\hat{k}_i(t) = k_i + \Pi_{\hat{p}_i \hat{\lambda}_i}(t - t_1 - R) + \sum_{j \in \Omega_i} \text{Ber}\left(e^{-\hat{\theta}_i(t_{i,j} + R - t_1)}\right),$$

où $\hat{p}_i = \hat{r} \exp(-\hat{\theta}_i R)$, $\hat{\lambda}_i$, $\hat{r}$ et $\hat{\theta}_i$ sont donnés par (2.13).

Notons $F_i(s)$ une transformée de Laplace de $\hat{\theta}_i$:

$$F_i(s) = \mathbb{E}[e^{-\hat{\theta}_i s}] = [1 + s/(\hat{\beta}_2 + T_i)]^{-\hat{\lambda}_2 - m_i}, \quad s \geq 0$$

Ainsi,

$$\begin{aligned} \mathbb{E}[\hat{k}_i(t) \mid data] &= k_i + (t - t_1 - R)\hat{r}F_i(R)\mathbb{E}[\hat{\lambda}_i] + \sum_{j \in \Omega_i} F_i(t_{i,j} + R - t_1), \\ \mathbb{V}[\hat{k}_i(t) \mid data] &= (t - t_1 - R)^2 \hat{r}^2 \mathbb{V}[e^{-\hat{\theta}_i R} \hat{\lambda}_i] + (t - t_1 - R)\hat{r} \mathbb{E}[e^{-\hat{\theta}_i R}] \mathbb{E}[\hat{\lambda}_i] \\ &\quad + \mathbb{V}\left[\sum_{j \in \Omega_i} \text{Ber}\left(e^{-\hat{\theta}_i(t_{i,j} + R - t_1)}\right)\right], \end{aligned}$$

où $\mathbb{E}[\hat{\lambda}_i]$ et $\mathbb{V}[\hat{\lambda}_i]$ sont donnés par (2.19), $\mathbb{V}[e^{-\hat{\theta}_i R} \hat{\lambda}_i] = F_i(R)\mathbb{E}[\lambda_i^2] - (F_i(R)\mathbb{E}[\hat{\lambda}_i])^2$, $\mathbb{E}[e^{-\hat{\theta}_i R}] = F_i(R)$. Finalement en notant $\Delta_{i,j} = R - t_1 - t_{i,j}$,

$$\begin{aligned} \mathbb{V}\left[\sum_{j \in \Omega_i} \text{Ber}\left(e^{-\hat{\theta}_i(t_{i,j} + R - t_1)}\right)\right] &= \sum_{j \in \Omega_i} F_i(\Delta_{i,j}) - F_i(2\Delta_{i,j}) \\ &\quad + \sum_{(j_1, j_2) \in \Omega_i^2} F_i(\Delta_{i,j_1} + \Delta_{i,j_2}) - F_i(\Delta_{i,j_1})F_i(\Delta_{i,j_2}). \end{aligned}$$

Modèle 5. Pour ce modèle $\hat{r}_i$ est également variable et donné par (2.16).

Notons $\hat{k}(t) = \sum_{i=1}^C \hat{k}_i(t)$ le nombre total de patients randomisés au temps $t > t_1 + R$. Pour chaque modèle, pour C assez grand ($C > 10$) on peut utiliser l'expression de l'espérance et de la variance de $\hat{k}(t)$ sachant les données pour créer les bornes prédictives de niveau $(1 - \delta)$ pour $\hat{k}(t)$ en utilisant une approximation normale comme dans [6]. Ces expressions sont donnés par la suite pour chaque modèle :

De manière similaire au modèle 4, on peut écrire

$$\begin{aligned} \mathbb{E}[\hat{k}_i(t) \mid data] &= k_i + (t - t_1 - R)\mathbb{E}[\hat{r}_i]F_i(R)\mathbb{E}[\hat{\lambda}_i] + \sum_{j \in \Omega_i} F_i(t_{i,j} + R - t_1), \\ \mathbb{V}[\hat{k}_i(t) \mid data] &= (t - t_1 - R)^2 \mathbb{V}[\hat{r}_i e^{-\hat{\theta}_i R} \hat{\lambda}_i] + (t - t_1 - R)\mathbb{E}[\hat{r}_i] \mathbb{E}[e^{-\hat{\theta}_i R}] \mathbb{E}[\hat{\lambda}_i] \\ &\quad + \mathbb{V}\left[\sum_{j \in \Omega_i} \text{Ber}\left(e^{-\hat{\theta}_i(t_{i,j} + R - t_1)}\right)\right], \end{aligned}$$

et utiliser les formules du modèle 4 précédentes pour le calcul des moyennes et variances des différentes variables de cette expression.

Remarque Il est également possible de considérer une distribution jointe d'un processus à deux composants $\{(\hat{n}_i(t), \hat{k}_i(t)), t \geq t_1 + R\}$. Sachant les données au temps t_1 , dans l'intervalle $[t_1, t_1 + R]$ ces processus sont indépendants, étant donné que pour $t \in [t_1, t_1 + R]$ le processus $\hat{k}_i$ ne dépend que des données avant t_1 . Pour $t > t_1 + R$, le processus $\{\hat{k}_i(t), t \geq t_1 + R\}$ peut être représenté par un amincissement du processus de Poisson $\{\hat{n}_i(t - R), t \geq t_1\}$ avec probabilité $\hat{r}$.

Prédiction de la proportion de patients randomisés.

Pour pouvoir approfondir l'analyse du phénomène des sorties d'étude, il est intéressant de considérer la proportion des patients randomisés par rapport au nombre de patients recrutés.

Notons $\{\hat{k}_i(t), t \geq t_1\}$ le processus des patients randomisés par le centre i . Étant donné que l'on sait que le nombre de patients recrutés $\hat{n}_i$ est un processus de Poisson-gamma, il suffit d'étudier la distribution jointe de $\hat{n}_i$ et la proportion de patients qui sont randomisés, i.e $\hat{k}_i/\hat{n}_i$. Rappelons qu'un patient randomisé au temps t a été recruté au temps $t - R$. Il est plus facile d'étudier le processus $\{\hat{k}_i(t)/\hat{n}_i(t - R), t \geq t_1 + R\}$, qui est approximativement le même que $\hat{k}_i/\hat{n}_i$ en supposant que R est petit comparé à la durée totale de la phase de recrutement (hypothèse réaliste au vu de la durée moyenne des phases de screening réelles qui varie de 2-4 semaines pour des périodes de recrutement de plusieurs années). Au lieu de considérer la distribution jointe du processus $\hat{n}_i$, nous considérons de manière équivalente la variable conditionnelle $\hat{k}_i(t)/\hat{n}_i(t - R)$ sachant $\tilde{n}_i(t - R) = N$ pour $t \geq t_1 + R$, où $\{\tilde{n}_i(t) = \hat{n}_i(t) - \hat{n}_i(t_1); t \geq t_1\}$ est le nombre de patients recrutés après t_1 .

Soit p_i la probabilité de randomisation du centre i et rappelons que conditionnellement à (r_i, θ_i) , nous avons $p_i = r_i$ dans les modèles 1-2 et $p_i = r_i e^{-\theta_i R}$ dans les modèles 3-5. Ainsi conditionnellement à $\tilde{n}_i(t - R) = N$, et sachant (r_i, θ_i) , le nombre de patients randomisés dans l'intervalle $[t_1 + R, t]$, $\hat{k}_i(t) - \hat{k}_i(t_1 + R)$, est une variable aléatoire binomiale $\text{Bin}(N, p_i)$ et est indépendant du nombre de patients screenés dans l'intervalle $[t_1, t_1 + R]$. Soit k_i le nombre total de patients randomisés jusqu'au temps t_1 . Il faut maintenant prendre en compte les patients randomisés dans l'intervalle $[t_1, t_1 + R]$.

Modèles 1-2. Soit ν_i le nombre de patients qui entre en période de screening dans le centre i dans l'intervalle $[t_1 - R, t_1]$. Ainsi, sachant ν_i et (r_i, θ_i) , le nombre de patients qui vont être randomisés dans l'intervalle $[t_1, t_1 + R]$ est une variable aléatoire binomiale $\text{Bin}(\nu_i, p_i)$ et est indépendant du nombre de patients screenés après $t_1 + R$. Donc, conditionnellement à $\hat{n}_i(t - R) = N$, nous avons la représentation

$$\hat{k}_i(t)/\hat{n}_i(t - R) = \frac{\text{Bin}(\nu_i + N, p_i) + k_i}{n_i + N}. \quad (2.23)$$

Étant donné que $p_i = r_i$, et en utilisant (2.6) où (2.8), il est possible de calculer la moyenne et la variance $\hat{k}_i(t)/\hat{n}_i(t - R)$ sachant les données au temps t_1 et $\tilde{n}_i(t - R) = N$.

Pour le modèle 1 nous avons $p_i \equiv \hat{r}$, ainsi

$$\begin{aligned} \mathbb{E}[\text{Bin}(\nu_i + N, \hat{r})] &= (\nu_i + N)\hat{r}, \\ \mathbb{V}[\text{Bin}(\nu_i + N, \hat{r})] &= (\nu_i + N)\hat{r}(1 - \hat{r}) \end{aligned}$$

Pour le modèle 2, $p_i \equiv r_i$ et $r_i \sim (\hat{\psi}_1 + k_i, \hat{\psi}_2 + n_i - k_i)$. Ainsi

$$\begin{aligned}\mathbb{E}[\text{Bin}(\nu_i + N, p_i)] &= (\nu_i + N)\mathbb{E}[r_i], \\ \mathbb{E}[(\text{Bin}(\nu_i + N, r))^2] &= (\nu_i + N)\mathbb{E}[r_i(1 - r_i)] + (\nu_i + N)^2\mathbb{E}[r_i^2]\end{aligned}$$

où

$$\mathbb{E}[r_i] = \frac{\hat{\psi}_1 + k_i}{\hat{\psi}_1 + \hat{\psi}_2 + n_i}, \quad \mathbb{E}[r_i^2] = \frac{(\hat{\psi}_1 + k_i)(\hat{\psi}_2 + n_i - k_i)}{(\hat{\psi}_1 + \hat{\psi}_2 + n_i)^2(\hat{\psi}_1 + \hat{\psi}_2 + n_i + 1)}$$

ce qui permet le calcul de la moyenne et de la variance de $\hat{k}_i(t)/\hat{n}_i(t - R)$ comme dans (2.23).

Modèles 3-5. Étant donné que nous avons plus d'informations disponibles au temps t_1 , les patients dont le résultat de la phase de screening correspond au cas $t_1 - R < t_{i,j} \leq t_1$ et $Y_{i,j} > t_1 - t_{i,j}$. Pour $1 \leq i \leq C$, rappelons que Ω_i est défini par (2.21) comme l'ensemble des indices et $\nu_i = \text{Card}(\Omega_i)$. Conditionnellement à θ_i , la probabilité pour ce patient d'être randomisé dans l'intervalle $[t_1, t_1 + R]$ est

$$\mathbb{P}[Y_{i,j} \geq R | Y_{i,j} > t_1 - t_{i,j}; \theta_i] = e^{-\theta_i(t_{i,j} + R - t_1)}$$

Sachant les données au temps t_1 et θ_i , le nombre de patients randomisés entre t_1 et $t_1 + R$ dans le centre i est la somme de ν_i variables aléatoires indépendantes de Bernoulli de probabilités $e^{-\theta_i(R - t_1 + t_{i,j})}$ noté $\text{Ber}(e^{-\theta_i(t_{i,j} + R - t_1)})$. Ainsi le proportion de patients randomisés au temps $t > t_1 + R$ sachant que N patients ont été recrutés dans l'intervalle $[t_1, t - R]$, les données à t_1 et les paramètres (r_i, θ_i) peuvent être représentés comme

$$\hat{k}_i(t)/\hat{n}_i(t - R) = \frac{\text{Bin}(N, p_i) + \sum_{j \in \Omega_i} \text{Ber}(e^{-\theta_i(t_{i,j} + R - t_1)}) + k_i}{n_i + N}, \quad (2.24)$$

où les différentes variables binomiales sont indépendantes.

Étant donné que $p_i = r_i e^{-\theta_i R}$ pour les modèles 3-5, et que (r_i, θ_i) sont donnés par (2.11), (2.13) où (2.16) selon le modèle considéré, nous pouvons calculer la moyenne et la variance de $\hat{k}_i(t)/\hat{n}_i(t - R)$ sachant les données au temps t_1 et $\hat{n}_i(t - R) = N$.

Pour le modèle 3, nous utilisons la représentation de (2.24), et nous avons $p_i \equiv r_i e^{-\theta_i R}$ où r_i et θ_i sont donnés par (2.11). Ainsi

$$\mathbb{E}[\text{Ber}(e^{-\theta_i(t_{i,j} + R - t_1)})] = \mathbb{E}[(\text{Ber}(e^{-\theta_i(t_{i,j} + R - t_1)}))^2] = e^{-\hat{\theta}(t_{i,j} + R - t_1)}$$

$$\mathbb{E}[\text{Bin}(N, p_i)] = N\hat{r}e^{-\hat{\theta}R}, \quad \mathbb{E}[(\text{Bin}(N, r))^2] = N\hat{r}e^{-\hat{\theta}R}(1 - e^{-\hat{\theta}R}) + N^2\hat{r}^2e^{-2\hat{\theta}R}$$

Pour le modèle 4,

$$\begin{aligned}\mathbb{E}[\text{Ber}(e^{-i(t_{i,j} + R - t_1)})] &= \mathbb{E}[(\text{Ber}(e^{-\theta_i(t_{i,j} + R - t_1)}))^2] \\ &= F(t_{i,j} + R - t_1; \hat{\alpha}_2 + m_i, \hat{\beta}_2 + T_i)\end{aligned}$$

où

$$F(s; a, b) = (1 - s/\beta)^{-\alpha} \quad (2.25)$$

est la transformé de Laplace d'une distribution $\Gamma(a, b)$.

$$\begin{aligned}\mathbb{E}[\text{Bin}(N, p_i)] &= N\hat{r}F(R; \hat{\alpha}_2 + m_i, \hat{\beta}_2 + T_i), \\ \mathbb{E}[(\text{Bin}(N, r))^2] &= N\hat{r}(F(R; \hat{\alpha}_2 + m_i, \hat{\beta}_2 + T_i) - \hat{r}F(2R; \hat{\alpha}_2 + m_i, \hat{\beta}_2 + T_i)) \\ &\quad + N^2\hat{r}^2F(2R; \hat{\alpha}_2 + m_i, \hat{\beta}_2 + T_i)\end{aligned}$$

Pour le modèle 5,

$$\begin{aligned}\mathbb{E}[\text{Ber}(e^{-\theta_i(t_{i,j}+R-t_1)})] &= \mathbb{E}[(\text{Ber}(e^{-\theta_i(t_{i,j}+R-t_1)}))^2] \\ &= F(t_{i,j} + R - t_1; \hat{\alpha}_2 + m_i, \hat{\beta}_2 + T_i)\end{aligned}$$

où F est donné par (2.25). De plus,

$$\begin{aligned}\mathbb{E}[\text{Bin}(N, p_i)] &= N\mathbb{E}[r_i]F(R; \hat{\alpha}_2 + m_i, \hat{\beta}_2 + T_i), \\ \mathbb{E}[\text{Bin}(N, r)^2] &= N\mathbb{E}[r_i]F(R; \hat{\alpha}_2 + m_i, \hat{\beta}_2 + T_i) - N\mathbb{E}[r_i^2]F(2R; \hat{\alpha}_2 + m_i, \hat{\beta}_2 + T_i) \\ &\quad + N^2\mathbb{E}[r_i^2]F(2R; \hat{\alpha}_2 + m_i, \hat{\beta}_2 + T_i)\end{aligned}$$

où

$$\mathbb{E}[r_i] = \frac{\hat{\psi}_1 + \tilde{k}_i}{\hat{\psi}_1 + \hat{\psi}_2 + n_i}, \quad \mathbb{E}[r_i^2] = \frac{(\hat{\psi}_1 + \tilde{k}_i)(\hat{\psi}_2 + n_i - \tilde{k}_i)}{(\hat{\psi}_1 + \hat{\psi}_2 + n_i)^2(\hat{\psi}_1 + \hat{\psi}_2 + n_i + 1)}$$

Finalement, il est possible de calculer la moyenne et la variance de la proportion globale de patients recrutés qui ont été randomisés au temps t sachant $\hat{n}_i(t - R) = N_i$ pour tout $1 \leq i \leq C$:

$$\hat{K}(t)/\hat{N}(t - r) = \left(\sum_{i=1}^M \hat{k}_i(t) \right) / \left(\sum_{i=1}^M \hat{n}_i(t - R) \right)$$

et nous obtenons l'intervalle de confiance en utilisant une approximation normale comme précédemment.

2.4 Prise en compte des interruptions de recrutement ²

L'effet des interruptions de recrutement sur la prédiction de la durée a déjà été présenté précédemment, il s'agit ici de voir si en ne considérant que cette perturbation, le modèle Poisson-gamma doit être ajusté en conséquence en ajoutant dans la modélisation une paramétrisation de ce phénomène, ou si le modèle peut dans une certaine mesure atténuer l'effet à lui seul.

Par la suite est présentée une paramétrisation spécifique de cette perturbation qui est alors incorporée au modèle, en supposant que les données d'interruptions soient disponibles, et également un mixte des deux modèles (le modèle de base Poisson-gamma et l'ajout des interruptions). Ces modèles seront comparés en terme d'erreur relative commise entre la prédiction de la durée du recrutement et la durée réelle.

Pour ce faire des études par simulations ont de nouveau été effectuées selon différentes hypothèses quant à la distribution des interruptions (distribution exponentielle, multinomiale et déterministe).

2. Ce travail a fait l'objet d'une soumission dans Contemporary Clinical Trials [49]

2.4.1 Le modèle Poisson-gamma avec incorporation des interruptions de recrutement

Le modèle incorporant les paramètres relatifs aux interruptions est présenté ici :

Supposons que le recrutement du centre i s'interrompt à l'instant noté $b_{i,j}$ pendant une période notée $d_{i,j}$.

Comme précédemment, fixons un temps intermédiaire t_1 . Les données collectées dans l'intervalle $[0, t_1]$ seront utilisées pour estimer les paramètres du modèle.

Pour le centre i , les données collectées sont :

- le nombre de patients recrutés jusqu'à t_1 noté $k_{1,i}$,
- le nombre d'interruptions jusqu'à t_1

$$j_{1,i} = \inf \{j : b_{1,i,j} < t_1, \text{ et } b_{1,i,j} + d_{1,i,j} \geq t_1\}$$

- $(b_{1,i,j}, j = 1, \dots, j_{1,i})$ et $(d_{1,i,j}, j = 1, \dots, j_{1,i})$ les temps d'interruptions et leurs durées jusqu'à t_1
- la durée d'activité jusqu'à t_1

$$\tau_{1,i} = (t_1 - u_i - \sum_{j=1}^{j_{1,i}} d_{1,i,j}) \vee 0$$

L'intensité pour le centre i est toujours gouvernée par λ_i qui est lui même distribué selon une loi $Gamma(\alpha, \alpha/\mu)$ qui est dépendant du temps et s'écrit pour un centre i comme :

$$t \rightarrow \lambda_i \chi[t \notin \mathcal{D}_{1,i}] \chi[t > u_i], \quad (2.26)$$

avec $\mathcal{D}_{1,i} = \{t; \exists j \in \llbracket 1, j_{1,i} \rrbracket; t \geq b_{i,j} \text{ et } t \leq b_{i,j} + d_{i,j}\}$.

Estimation

Pour estimer (α, μ) le théorème suivant assure que la même stratégie est appliquée que celle présentée en Section 1.3.2, seul la définition des τ_i diffère.

Théorème 4 *L'estimation du maximum de vraisemblance $(\hat{\alpha}, \hat{\mu})$ des paramètres (α, μ) est obtenue par maximisation de la fonction :*

$$M_C^{Gamma}(\alpha, \beta) = \alpha \ln(\beta) - \ln \Gamma(\alpha) + \frac{1}{C} \sum_{i=1}^C [(\alpha + k_{1,i}) - (\alpha + k_{1,i}) \ln(\beta + \tau_{1,i})].$$

Preuve Comme montré dans [46], lorsqu'on utilise une intensité λ non déterministe, la fonction de densité d'un processus de Poisson f est remplacée par $\mathbb{E}_\lambda[f]$.

Sachant cela :

$$\begin{aligned}
 \mathbb{P}[N_i(t_1) = n_i] &= \mathbb{E}_{p_{\alpha,\beta}} [\mathbb{P}[N_i(t_1) = n_i]] \\
 &= \mathbb{E}_{p_{\alpha,\beta}} \left[\exp\left(-\int_0^{t_1} A(\lambda_i, s) ds\right) \frac{(\int_0^{t_1} A(\lambda_i, s) ds)^{n_i}}{n_i!} \right] \\
 &= \mathbb{E}_{p_{\alpha,\beta}} \left[\exp\left(-\lambda_i(t_1 - u_i - \sum_{j \in C^i} d_{i,j})\right) \frac{(\lambda_i(t_1 - u_i - \sum_{j \in C^i} d_{i,j}))^{n_i}}{n_i!} \right] \\
 &= \mathbb{E}_{p_{\alpha,\beta}} \left[\exp(-\lambda_i \tau_i) \frac{(\lambda_i \tau_i)^{n_i}}{n_i!} \right] \\
 &= \frac{\tau_i^{n_i}}{n_i!} \int_0^\infty \exp(-x \tau_i) x^{n_i} p_{\alpha,\beta}(x) dx
 \end{aligned}$$

Par indépendance entre les processus $\mathcal{N}_i$ pour $i = 1 : C$, on obtient :

$$\begin{aligned}
 \mathbb{P}[N_i(t_1) = n_i] &= \prod_{i=1}^C \frac{\tau_i^{n_i}}{n_i!} \int_0^\infty \exp(-x \tau_i) x^{n_i} p_{\alpha,\beta}(x) dx \\
 &= \prod_{i=1}^C \frac{\Gamma(\alpha + n_i)}{n_i! \Gamma(\alpha)} \frac{\beta^\alpha \tau_i^{n_i}}{(\beta + \tau_i)^{\alpha + n_i}}
 \end{aligned}$$

Le résultat est obtenu en sachant que la fonction log-vraisemblance est égale au log de la précédente équation [10].

Prédiction

Lorsque l'on considère des interruptions de recrutement la difficulté se trouve dans la prédiction du processus de Poisson global non-homogène. Pour surmonter cette difficulté, le processus est séparé en deux parties, comme montré sur la Figure 2.11.

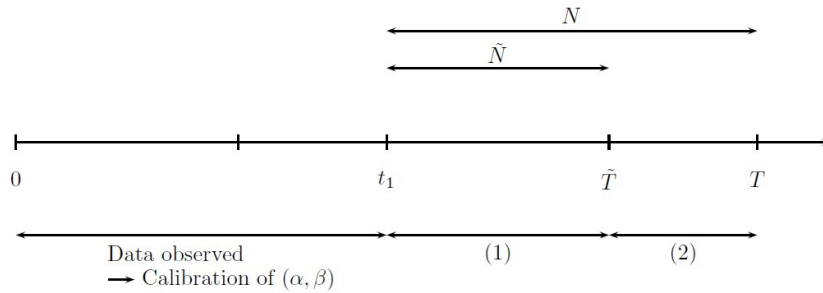


FIGURE 2.11 – Processus impliqués dans la prédiction de T

Différents processus sont impliqués dans cette partie. Premièrement le processus N qui est non-homogène pour lequel les centres ont des intensités de la même forme que dans l'expression 2.26, mais cette fois les interruptions prises en compte sont celles survenant après le temps t_1 , dont aucune information n'est disponible sauf l'estimation de

sa distribution avant l'instant t_1 . Ce processus est le processus réel décrivant la dynamique de recrutement étudiée, processus qui est approximé comme expliqué par la suite.

Le paramètre d'intérêt pour ce processus est $T = \{\inf_{t \geq 0} : N(t) = n\}$.

Deuxièmement le processus $\tilde{N}$ qui est un processus de Cox homogène (qui commence en t_1) dont l'intensité est donnée par :

$$t \rightarrow \sum_{i=1}^C \hat{\lambda}_i \chi[t \geq t_1] , \quad (2.27)$$

Le paramètre d'intérêt pour ce processus est $\tilde{T} = \{\inf_{t \geq 0} : \tilde{N}(t) = n\}$ c'est à dire la durée (1).

Selon sa définition le processus $\tilde{N}$ ne prend pas en compte les interruptions, il est donc logique que :

$$\tilde{T} \leq T$$

Finalement la durée des interruptions après t_1 (2) est approché selon la méthode présenté ci-dessous.

Prédiction de la durée de recrutement

Pour approcher le processus réel N , il est séparé en deux processus, dont l'un est un processus de Poisson $\tilde{N}$ amenant à la prédiction $\mathbb{E}[\tilde{T}]$ comme effectué par le modèle standard, et un processus d'interruptions lui aussi approché.

Notons que

$$\mathbb{E}[T] = \mathbb{E}[\tilde{T}] + \mathbb{E}[T - \tilde{T}]. \quad (2.28)$$

$\mathbb{E}[\tilde{T}]$ est relié à un processus de Poisson homogène ainsi le Théorème 2 s'applique et nous donne une estimation de cette durée.

$\mathbb{E}[T - \tilde{T}]$ est estimé en supposant que la durée totale des interruptions est distribuée uniformément sur $[0, T^r]$, ainsi les données sur $[0, t_1]$ servent d'approximation de la durée des interruptions après t_1 selon la formule :

$$BC_1 = \frac{\mathbb{E}[\tilde{T}]}{\sum_{i=1}^C (t_1 - u_i)} \sum_{i=1}^C \sum_{j=1}^{j_{1,i}} d_{i,j}$$

Il faut noter que l'estimation de $\mathbb{E}[T - \tilde{T}]$ proposée est surement biaisée. Une étude par simulation permet de quantifier ce biais et sur la nécessité ou non de prendre en compte les interruptions de recrutement lorsque le modèle Poisson-gamma est appliqué à une phase de recrutement.

Stratégie alternative

Cette partie montre une méthode qui est un 'mixte' de l'approche standard et de l'approche précédente.

Le but est de considérer seulement les interruptions dont la durée est supérieure à un seuil d prédéfini.

Après avoir fixé le temps t_1 , les données collectées entre $[0, t_1]$ sont utilisées pour calibrer les paramètres du modèle.

Pour le centre i , les données collectées sont :

- le nombre de patients recrutés jusqu'à t_1 noté $k_{1,i}$,
- le nombre d'interruptions dont la durée est supérieure à d jusqu'à t_1

$$j_{1,i} = \inf\{j : b_{1,i,j} < t_1, \text{ et } b_{1,i,j} + d_{1,i,j} \geq t_1 \text{ et } d_{i,j} \geq d\}$$

- $(b_{1,i,j}, j \in J$ et $(d_{1,i,j}, j \in J$, respectivement les instants et durées des interruptions jusqu'à t_1 supérieure à d , avec

$$J = \{j = 1, \dots, j_{1,i} : d_{1,i,j} \geq d\}$$

- la durée d'activité jusqu'à t_1

$$\tau_{1,i} = (t_1 - u_i - \sum_{j \in J} d_{1,i,j}) \vee 0$$

L'intensité pour le centre i est toujours gouvernée par λ_i qui est lui même distribué selon une loi $Gamma(\alpha, \alpha/\mu)$ qui est dépendant du temps et s'écrit pour un centre i comme :

$$t \rightarrow \lambda_i \chi[t \notin \mathcal{D}_{1,i}] \chi[t > u_i], \quad (2.29)$$

avec $\mathcal{D}_{1,i} = \{t; \exists j \in J; t \geq b_{i,j} \text{ et } t \leq b_{i,j} + d_{i,j}\}$.

Avec ces notations, et définitions, l'estimation des paramètres ainsi que la prédiction de la durée du recrutement se font comme en Section 2.4.1.

2.4.2 Procédure de génération des données

Considérons un essai multicentrique impliquant $C = 60$ centres. Le nombre NSN sera $n = 720$ patients dont on estime le temps de recrutement à 365 jours. Dans l'optique d'étudier et de comparer les différentes approches présentées précédemment à l'approche standard, plusieurs scénarios de générations d'interruptions sont considérés, ces scénarios sont :

- **Scénario 1 : génération Exponentielle.** Les instants et durées des interruptions sont générés selon des lois exponentielles.

Les temps d'interruptions sont générés selon une loi exponentielle de paramètre $\frac{1}{60}$ (une interruption tout les 2 mois), et les durées des interruptions sont elles générées selon une loi exponentielle de paramètre $\frac{1}{14}$ (durée de 2 semaines)

- **Scenario 2 : génération Multinomiale.** Les instants et durées des interruptions sont générés selon des lois multinomiales.
Les temps d'interruptions sont générés selon une loi exponentielle de paramètre $\frac{1}{60}$ (une interruption tout les 2 mois), et les durées des interruptions sont elles générées selon une loi multinomiale à 5 niveaux (2, 4, 8, 16 et 32 jours) dont le vecteur de probabilité est créé de sorte que la durée de chaque niveau multipliée par leur nombre d'occurrence soit la même pour tout les niveaux (le niveau 2 jours aura une probabilité deux fois plus importante d'arriver que le niveau 4 jours)
- **Scenario 3 : génération Déterministe.** Les instants et durées des interruptions sont fixes : deux jours par semaine (week-end) et une semaine tout les deux mois (congés).

Pour étudier les performances prédictives des différentes approches, des études par simulation sont effectuées.

Les dynamiques de recrutement sont générées en incorporant les interruptions selon le scénario considéré.

La dynamique est connue dans sa totalité ainsi la durée réelle notée T_0 est connue. L'étude consiste à utiliser les données jusqu'à t_1 et utiliser les résultats de l'approche classique et de l'approche présentée en section 2.4.1 pour estimer la durée du recrutement de l'essai. Pour répondre à la question "Est il nécessaire de prendre en compte les interruptions de recrutement dans le modèle Poisson-gamma", 3 stratégies sont considérées et comparées :

- **Stratégie 1 : ne pas prendre en compte les interruptions.** Les instants et durées des interruptions de recrutement ne sont pas collectées. Les paramètres du modèle Poisson-gamma ainsi que la prédiction de la durée de recrutement sont obtenues de manière standard. La durée prédite est notée T_1 .
- **Stratégie 2 : Prise en compte de toutes les interruptions.** Les instants et durées des interruptions de recrutement sont collectés. Cela permet l'utilisation du modèle présenté en Section 2.4.1. La durée prédite est notée T_2 .
- **Stratégie 3 : Prise en compte des grandes interruptions.** Seules les interruptions dont la durée est plus grande que $d = 14$ jours sont collectés et utilisées tel que présenté en Section 2.4.1. La durée prédite est notée T_3 .

Dans un soucis de simplicité, tout les centres sont initialisés à l'instant $t = 0$ (les dates d'ouvertures sont toutes égales à 0).

La génération des données se découpe en deux phases :

Étape 1 : La génération de $R = 1000$ processus de recrutement globaux

$\{N^r(t), 0 \leq t \leq T_0^r, 1 \leq r \leq R\}$ où T_0^r dénote le premier temps vérifiant $N^r(t) = 720$. La procédure de génération est la suivante :

1. Génération des interruptions selon le scénario considéré sur une période de 730 jours (deux fois le temps attendu sans interruptions).
2. Génération des intensités selon une distribution $Gamma(2, 60.8)$. Les paramètres (2, 60.8) de la loi gamma sont ceux choisis dans [10] pour s'assurer du côté réaliste du processus de recrutement.

3. Considération de la fonction d'intensité perturbée telle que exprimée par l'équation 2.26.
4. Génération des intensités de recrutement jusqu'au temps $t = 730$ jours.
5. Identification de T_0^r et raccourcissement du processus sur $[0, T_0^r]$.

Étape 2 : Sachant un temps intermédiaire $t_1 = 182$ jours. Pour chaque simulation $r = 1, \dots, R$ et chaque stratégie considérée $s = 1, 2, 3$,

1. Estimation de (α_s^r, μ_s^r) de manière standard pour $s = 1$ ou par application du théorème 4 pour $s = 2$ et $s = 3$ à partir des données collectées sur l'intervalle $[0, t_1]$.
2. Calcul de la prédiction de la durée de recrutement T_s^r de manière standard pour $s = 1$ ou par application de l'expression 2.28 pour $s = 2$ et $s = 3$.

Les performances de l'application du modèle au temps t_1 sont mesurées grâce à l'erreur absolue définie par :

$$E_{s,s'} = \frac{1}{R} \sum_{r=1}^R |T_s^r - T_{s'}^r|, \quad \text{for } s = 0, 1, 2, 3, s' = 0, 1, 2, 3 \text{ et } s \neq s'. \quad (2.30)$$

Remarque La valeur seuil de la stratégie 3 (14 jours) ne considère aucune des interruptions du scénario 3 (de durée maximum de 7 jours). Les résultats de ce scénario selon la stratégie 3 seront donc égaux à ceux obtenus avec la stratégie 1.

2.4.3 Résultats et discussion

Résultats sur le biais

Le premier point investigué lors de cette étude est le biais inhérent à chaque stratégie d'estimation et de prédiction. Les résultats sont collectés dans le Tableau 2.4. Pour chaque scénario sont identifiés la moyenne (sur l'ensemble des simulations conduites) des durées des dynamique de recrutement simulées de paire avec les intervalles de confiance à 95 % (correspondant aux 25-ème et 975-ème valeurs de l'échantillon ordonné). De plus, pour la stratégie p ($p = 1, 2, 3$), la moyenne et l'intervalle de confiance à 95 % sont calculés.

Le biais est mesuré par la moyenne des erreurs absolues $E_{0,p}$ entre la prédiction et la durée réelle de chaque simulation. Le tableau est également enrichi par les valeurs $S_{p,0}$ mesurant la proportion de sur-estimation de la prédiction de la stratégie p sur l'ensemble des simulations. S'ajoutent également les Figures 2.12-2.14. Figure 2.12 (resp. 2.13, 2.14) réfère au scénario 1 (resp. 2, 3). Chaque figure est composée des histogrammes des prédictions des toutes les simulations ($R = 1000$) pour chacune des stratégies ainsi que les fonctions de densités estimées de chaque stratégie auxquelles s'ajoute celle de la vraie durée. Il faut garder à l'esprit que selon la Remarque 2.4.2, la stratégie 3 n'est pas étudiée pour le scénario 3.

TABLEAU 2.4 – Durée moyenne comme fonction de la stratégie ainsi que l'intervalle de confiance (IvC) à 95 %. L'erreur absolue entre la prédiction et la durée réelle ($E_{0,p}$) comme fonction de la stratégie ainsi que l'intervalle de confiance à 95% et la proportion de sur-estimation ($S_{p,0}$).

		Durée réelle	Stratégie 1	Stratégie 2	Stratégie 3
Scénario 1	Moyenne	384.37	381.92	374.66	377.73
	IvC (95 %)	[351,415]	[340,431]	[336,421]	[337,426]
	$E_{0,p}$		13.81	15.42	14.41
	IvC (95 %)		[2,13]	[1,40]	[0,38]
	$S_{p,0}$		0.41	0.27	0.32
Scénario 2	Moyenne	337.92	337.8	336.41	337.16
	IvC (95 %)	[303,373]	[304,374]	[303,372]	[353,406]
	$E_{0,p}$		9.78	9.88	9.82
	IvC (95 %)		[0,29]	[1,29]	[0,29]
	$S_{p,0}$		0.44	0.40	0.42
Scénario 3	Moyenne	369.17	367.95	356.94	-
	IvC (95 %)	[343,398]	[330,412]	[321,398]	-
	$E_{0,p}$		11.04	14.88	-
	IvC (95 %)		[1,32]	[1,38]	-
	$S_{p,0}$		0.44	0.16	-

La première observation est que par rapport à la moyenne des durées réelles (troisième colonne), la meilleure stratégie est la première pour tous les scénarios, c'est à dire la non considération des interruptions. C'est à dire que en comparaison des deux autres approches, le modèle Poisson-gamma à l'air de mieux incorporer les interruptions que les deux autres approches.

Deuxièmement en terme de moyenne des erreurs absolues c'est de nouveau la stratégie 1 qui est la meilleure pour tous les scénarios car la plus faible (variabilité autour de la durée réelle moindre).

Troisièmement toutes les stratégies sous-évaluent en moyenne la durée réelle, la meilleure étant de nouveau la première stratégie (dont la proportion est la plus proche de 0.5).

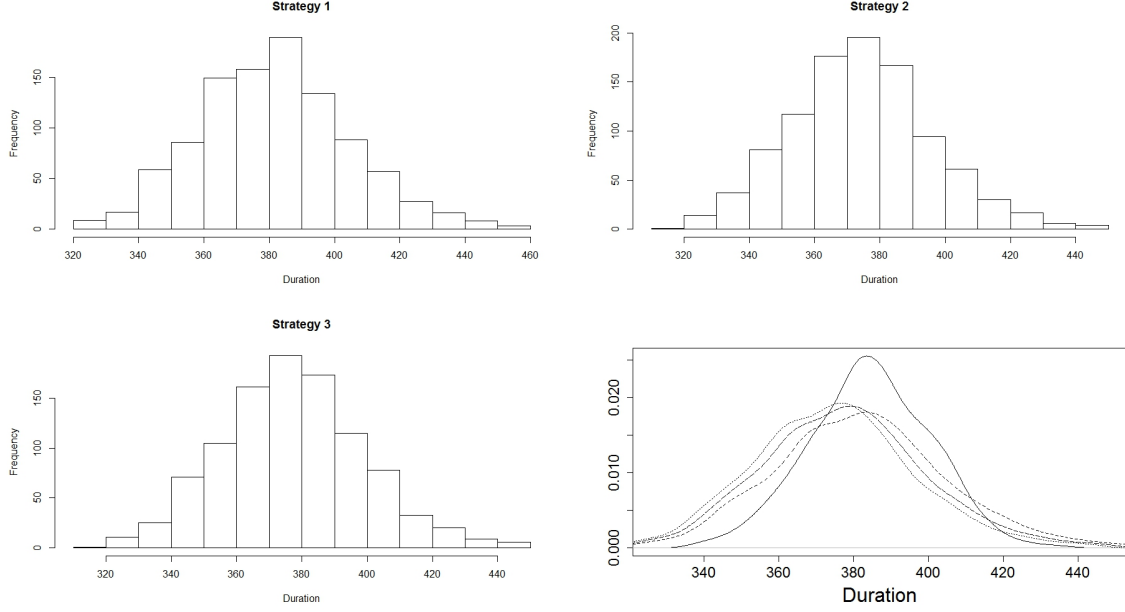


FIGURE 2.12 – Histogrammes des durées prédites pour tous les scénarios et comparaison des densités approchées : interruptions exponentielle (ligne pleine : T_r , trait : T_1 , points : T_2 , trait long : T_3)

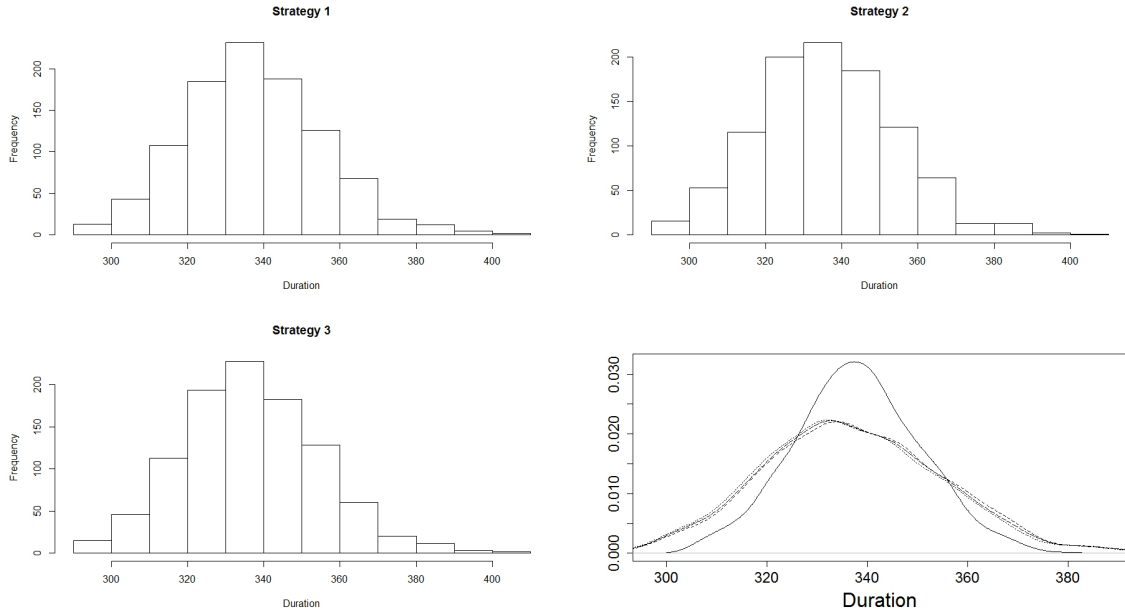


FIGURE 2.13 – Histogrammes des durées prédites pour tous les scénarios et comparaison des densités approchées : interruptions multinomiales (ligne pleine : T_r , trait : T_1 , points : T_2 , trait long : T_3)

Comparaison des stratégies

Les Tableaux 2.5-2.7 collectent les résultats des études par simulation pour chaque scénario. Le Tableau 2.5 réfère au scénario 1, le Tableau 2.6 au scénario 2 et le Tableau 2.7 au scénario 3. Chaque bloc des tableaux correspondent aux comparaison entre deux stratégies. Ces comparaisons entre une stratégie s et une autre s' sont données en terme d'erreur absolue $E_{s,s'}$ ainsi que les intervalles de confiance à 95% et les proportion $S_{s,s'}$

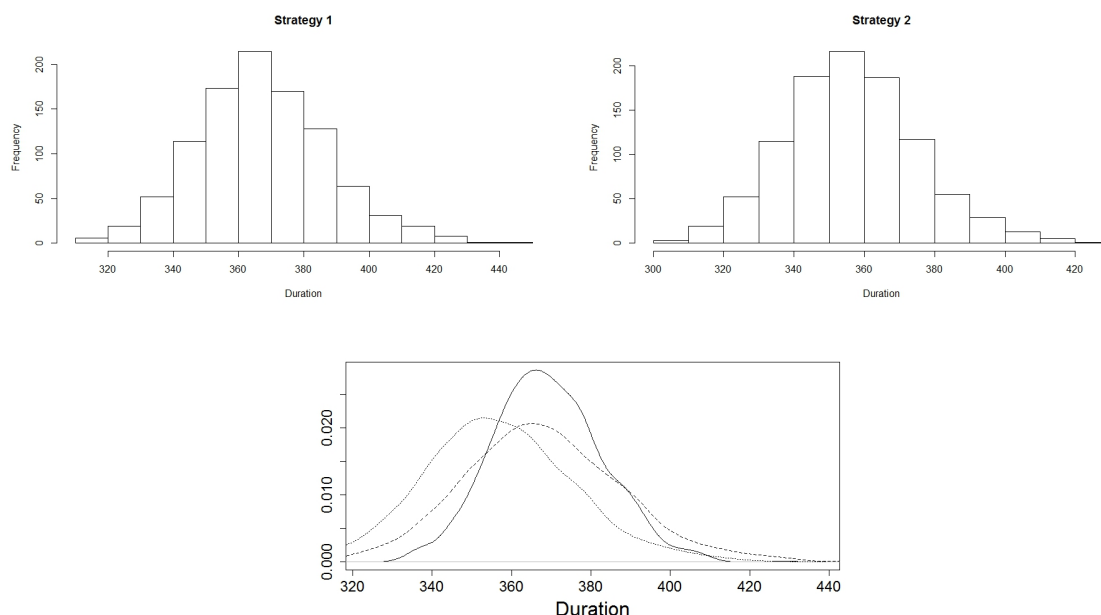


FIGURE 2.14 – Histogrammes des durées prédites pour tous les scénarios et comparaison des densités approchées : interruptions déterministes (ligne pleine : T_r , trait : T_1 , points : T_2)

de sur-estimations, c'est à dire la proportion du nombre de fois où la prédiction de la stratégie s est plus grande que celle de la stratégie s' .

TABLEAU 2.5 – Résultats du scénario 1. Comparaison des stratégies en terme d'erreur absolue ainsi que les intervalles de confiance à 95% et les proportions de sur-estimations.

		Stratégie 1	Stratégie 2
Stratégie 2	$E_{2,p}$	7.27	-
	IvC (95 %)	[2,13]	-
	$S_{p,2}$	0.99	-
Stratégie 3	$E_{3,p}$	4.25	3.07
	IvC (95 %)	[0,10]	[1,5]
	$S_{p,3}$	0.94	0.01

TABLEAU 2.6 – Résultats du scénario 2. Comparaison des stratégies en terme d'erreur absolue ainsi que les intervalles de confiance à 95% et les proportions de sur-estimations.

		Stratégie 1	Stratégie 2
Stratégie 2	$E_{2,p}$	1.44	
	IvC (95 %)	[0,3]	
	$S_{p,2}$	0.82	
Stratégie 3	$E_{3,p}$	0.82	0.83
	IvC (95 %)	[0,2]	[0,2]
	$S_{p,3}$	0.55	0.04

Concernant l'étude du Tableau 2.5, il est visible que les stratégies sont assez similaires. La plus grande différence est en effet de 7 jours (sur une période de 180

TABLEAU 2.7 – Résultats du scénario 3. Comparaison des stratégies en terme d'erreur absolue ainsi que les intervalles de confiance à 95% et les proportions de sur-estimations.

Stratégie 2	Stratégie 1	
	$E_{2,1}$	11.01
	IvC (95 %)	[9,14]
	$S_{1,2}$	1

de prédiction ($T^r - t_1$)). Les intervalles de confiance ne sont pas très grand et leurs grandeurs sont de nouveau similaires. Finalement les proportions de sur-estimation montrent que la stratégie 1 est à 99% plus grande que la stratégie 2 et de 94% par rapport à la stratégie 3.

Pour le Scénario 2, le Tableau 2.6 indique que les résultats obtenus sont équivalents à ceux obtenus en considérant le scénario 1, que ce soit en terme de proximité des moyennes et des fortes proportion de sur-estimation de la stratégie 1 par rapport aux deux autres.

Finalement, Scénario 3, le Tableau 2.7 indique que l'écart des erreurs absolues est de nouveau comparable avec les deux précédents scénarios, ainsi que la sur-estimation de la stratégie 1 sur la seconde.

Performance prédictive des stratégies

Le comportement prédictif d'une stratégie est un concept individuel au contraire de l'étude précédente où le comportement moyen était analysé. La qualité prédictive des différentes stratégies est étudiée par comparaison des 1000 prédictions effectuées à la durée réelle. Ceci se représente graphiquement par comparaison du nuage de points formé par les 1000 couples formés en abscisse par les prédictions et en ordonnée par les durées réelles de chacune des générations à la droite $y = x$ représentant des prédictions parfaites (exactement égales aux durées réelles). Ces représentations graphiques sont montrés en Figure 2.15-2.17. La Figure 2.15 correspond au Scénario 1, la Figure 2.16 au Scénario 2 et la Figure 2.17 au Scénario 3. A partir de ces nuages de points, nous effectuons des régressions linéaires pour étudier la proximité des prédictions aux durées réelles.

Pour chaque figure, la ligne continue représente la droite $y = x$ et la ligne pointillé la droite de régression. Toutes les régressions aboutissent à des modèles prédictives pertinent, les statistiques F sont toutes supérieures à 740 avec des p-values inférieure à 0.01. Aucune auto-correlation n'est observée mais l'homoscedasticité et la normalité des résidus sont confirmées.

La Figure 2.15 montre comme l'étude précédente que la stratégie 1 est la plus proche de la durée réelle, ainsi la pente de la droite de régression est de 0.94, contre 0.87 pour la stratégie 2 et 0.89 pour la 3 (la conformité des pentes par rapport à celle de la droite $y = x$ est toujours vérifiée grâce à des test de Wald). Nous observons également le phénomène de sous-évaluation des stratégies par rapport aux durées réelles (les droites de régression sont sous la droite $y = x$). Finalement aucun biais du scénario n'est observé car aucune différence constante n'est observable entre les droites de régressions et la droite $y = x$.

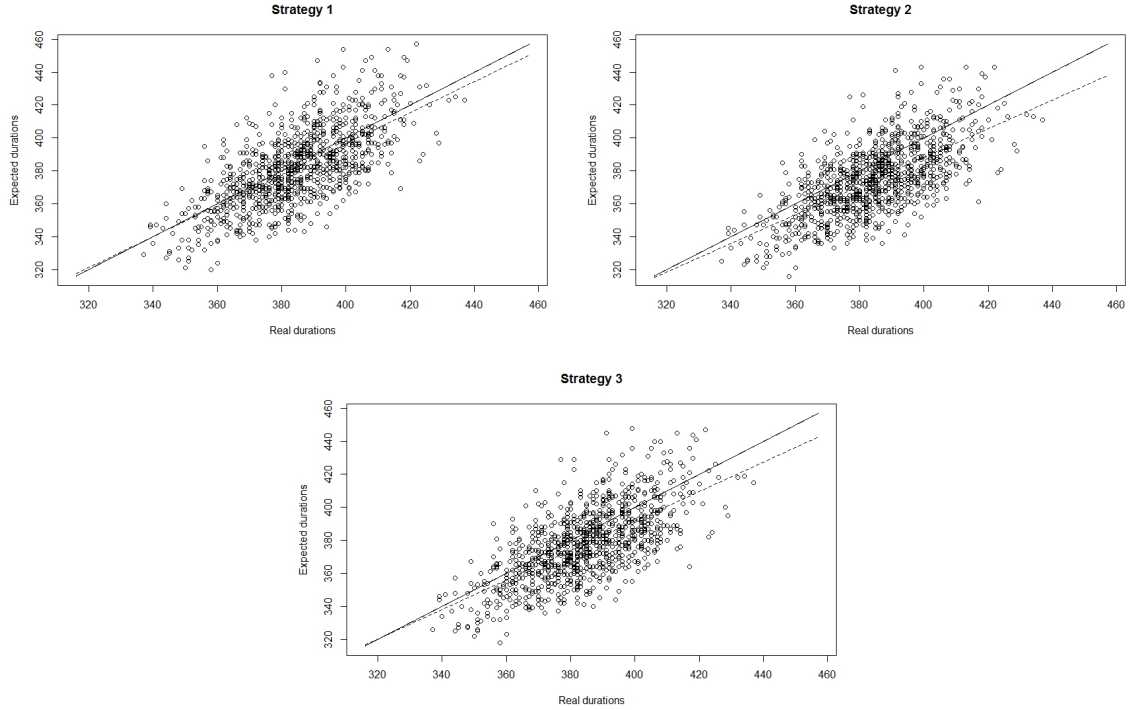


FIGURE 2.15 – Droite de régression entre la durée prédite et la durée réelle pour chaque stratégie : scénario exponentiel (ligne droite : droite $y = x$, ligne pointillée : droite de régression)

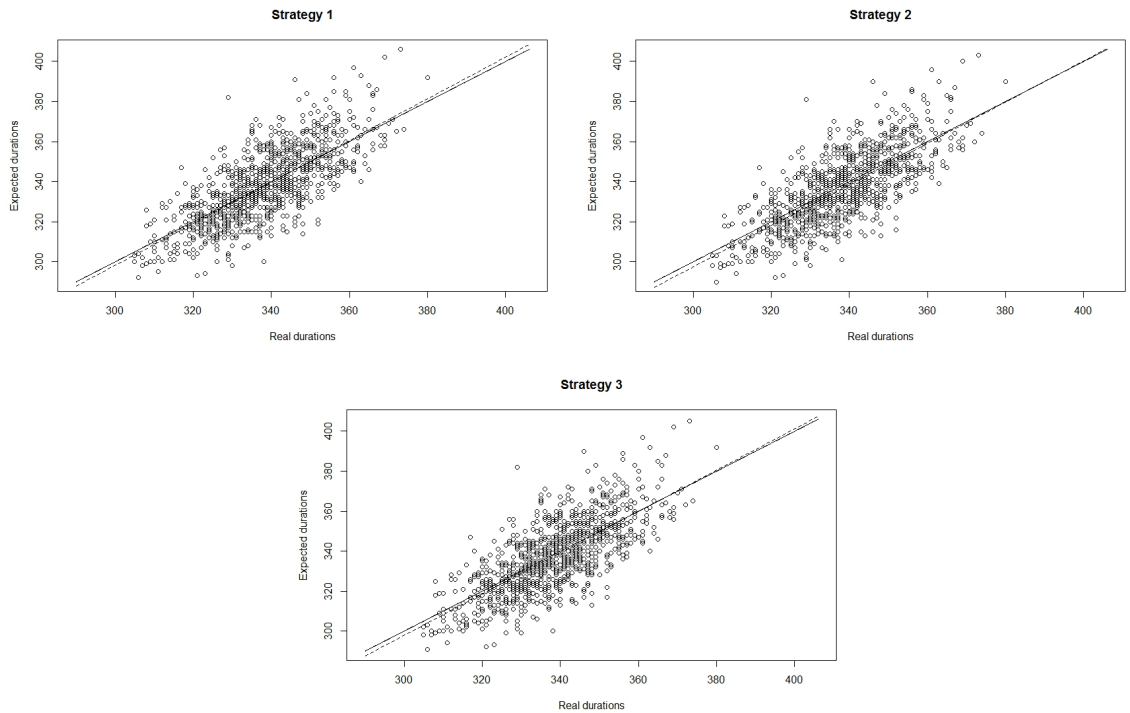


FIGURE 2.16 – Droite de régression entre la durée prédite et la durée réelle pour chaque stratégie : scénario multinomial (ligne droite : droite $y = x$, ligne pointillée : droite de régression)

La Figure 2.16 montre une situation plus complexe, en effet il est difficile de distinguer la meilleure stratégie lorsque le Scénario 2 est considéré. Les pentes de chacune des droites de régressions sont très proches (1.04 pour la stratégie 1, 1.03 pour les stratégies

2 et 3). Aucune sous ou sur-estimation n'est visible, ainsi qu'aucun biais inhérent au scénario 2. Cette fois ci seule l'étude des Tableaux 2.4 et 2.6 permet une comparaison d'efficacité entre les stratégies.

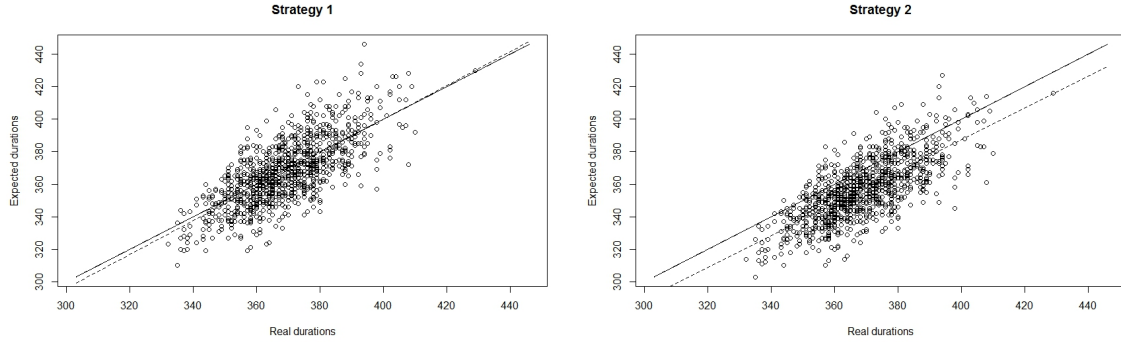


FIGURE 2.17 – Droite de régression entre la durée prédite et la durée réelle pour chaque stratégie : scénario déterministe (ligne droite : droite $y = x$, ligne pointillée : droite de régression)

Dans le cas du scénario 3, de nouveau la stratégie 1 est la plus performante en terme de proximité de la prédiction à la durée réelle. Les pentes sont assez proches entre les deux stratégies (1.04 pour la première contre 0.98 pour la seconde), mais l'ordonnée à l'origine de la première stratégie est plus proche que la seconde. Cette écart constant entre la droite de régression pour la seconde stratégie avec la droite $y = x$ nous informe d'un biais dans l'approximation faite BC_1 . Finalement la stratégie 2 sous-estime la durée réelle, tandis que ni sous ou sur-estimation n'est visible pour la stratégie 1.

2.4.4 Conclusions

Le but de cette étude était d'observer l'efficacité du modèle Poisson-gamma lorsque l'une de ses hypothèses n'était pas vérifiée, à savoir la constance des intensités de recrutement au cours du temps, et d'étudier si la prise en compte de ces interruptions est une meilleure option à choisir plutôt que de laisser le modèle les incorporer de lui même dans ces estimations de paramètres.

Nous pouvons ainsi dire que dans le cas d'interruptions dont la génération est exponentielle, la prise en compte des breaks est une perte de temps sans gain de performance prédictive.

Sachant que chaque stratégie (pour toutes les scénarios) a tendance à sous-évaluer la durée réelle, la stratégie 1 est la plus efficace (car elle donne presque toujours des prédictions les plus grandes par rapport aux deux autres). Il est donc plus efficace de laisser le modèle incorporer ces interruptions plutôt que de les ajouter comme proposer dans ce paragraphe.

Nous avons de plus vu que les prédictions faites par chaque stratégie étaient proches des durées réelles (erreur ne dépassant pas 15 jours sur des recrutement de 365 jours), ce qui est très motivant.

Chapitre 3

Applications du modèle

Outre l'utilisation initiale du modèle Poisson-gamma pour la prédiction du temps réelle que met une dynamique de recrutement multicentrique pour atteindre un nombre de patients cible n , le modèle peut en outre informer ou modéliser de multiples aspects de cette dynamique. Cette Section se veut un état de l'art de multiples applications possibles du modèle Poisson-gamma testées ou modélisées depuis sa présentation par Anisimov et al. [9].

3.1 A l'étude de faisabilité d'un essai

L'étude de la faisabilité d'un essai est une question difficile mais fondamentale. En effet il est fondamental d'estimer la durée d'un essai clinique et donc la durée de la phase de recrutement avant même son commencement puisqu'en général c'est la composante aléatoire de la question. Difficile car cela consiste en l'estimation de paramètres (notamment les intensités de recrutement) sans données sur la dynamique de recrutement elle même. Usuellement la faisabilité se base sur les valeurs données par les investigateurs de chaque centre qui ont une idée (plus ou moins) précise de leur potentiel de recrutement. Cependant il est tentant de se demander si l'utilisation de données d'essais terminés, qu'elles soient sélectionnées ou non selon des critères de choix prédéfinis, pourrait-être utile.

3.1.1 Indicateurs selon les paramètres a priori d'experts

Dès lors qu'une phase de recrutement de phase III est décidé, une approximation de sa durée est effectuée sachant donné par les investigateurs la valeur des paramètres du modèle prédictif.

Le manque de données est alors surpassé si l'on considère que ces paramètres sont pertinents selon l'étude qui va être conduite. Seulement cette question est encore très difficile à évaluer tant les cofacteurs sont multiples et non injectés dans l'estimation des paramètres.

L'idée d'utiliser des données historiques a été étudiée dans [13] selon l'hypothèse que l'estimation des paramètres se fait sur un grand nombre d'études (chacune incluant une phase de recrutement) où méta-analyse.

La conclusion qui en ressort dénote du côté problématique de l'estimation de ces paramètres qui ne saurait en général représenter la totalité des recrutements effectués (l'idée de construire un estimateur global pour toute phase de recrutement est très complexe et nécessiterait d'énormes quantités de données : pas seulement sur les dynamiques de recrutement mais comme introduit précédemment des cofacteurs également).

L'idée d'estimer les valeurs des paramètres sur des études historiques sans ajout de cofacteurs influents où sans hypothèses de 'proximité' entre les études est pour le moment insuffisante [13]. Plusieurs manières de considérer cette proximité peuvent être questionnées et étudiées, par exemple dans le cas de centres investigateurs sollicités pour plusieurs études sur la 'même' population cible.

Cette question a été investiguée dans un article en cours de rédaction [47] se basant sur des données réelles de recrutement. Pour mettre en place une telle méthodologie, 3 études dont les protocoles sont très proches et dont les centres investigateurs sont pour certains les mêmes ont été utilisées. L'idée est d'étudier la validité du critère de choix qui est émis, à savoir l'utilisation de données historiques (données récoltées sur des dynamiques de recrutement finies et où les données accessibles satisfont les hypothèses du modèle Poisson-gamma) est pertinent dès lors qu'une grande proportion des patients recrutés entre les deux études comparées l'ont été par des centres communs à ces deux études (l'étude historique : qui est finie, et l'étude en cours de mise en place : qui n'est pas encore commencée). Il s'agit ici des études IFM 2005, IFM 2009 et IFM 2014. La description des centres impliqués dans ces études est résumée dans le tableau 3.1.

L'hypothèse selon laquelle une grande proportion de patients est recrutée par un même ensemble de centres est difficilement vérifiable en pratique car les données de recrutement de la nouvelle étude (dont le recrutement n'est pas encore commencé) ne sont pas encore accessibles (à l'exception des centres inclus). C'est pourquoi l'utilisation des deux études historiques donc finies (études IFM 2005 et IFM 2009) sont d'abord utilisées pour appuyer expérimentalement cette hypothèse et comparer différents choix portés sur les valeurs des intensités des dynamiques de recrutement considérées.

TABLEAU 3.1 – Description des processus de recrutement des essais IFM 2005, IFM 2009 et IFM 2014.

	IFM 2005	IFM 2009	IFM 2014
Durée de l'essai (attendue pour 2014) en jour	914	746	730
Nombre total de centres	76	67	47
Nombre de centres partagés avec 2005	-	50(74%)	36(77%)
Nombre de centres partagés avec 2009	50(66%)	-	41(87%)
Nombre total de patients	611	693	-
Recrutés par les centres partagés avec 2005	-	562 (81%)	-
Recrutés par les centres partagés avec 2009	543 (89%)	-	-

L'étude du tableau 3.1 montre ainsi une 'grande' proportion de centres communs d'une part (en considérant l'étude IFM 2005 comme référent, 66% des centres participant à cette étude sont partagés avec l'étude IFM 2009 et cette proportion passe à

87% avec l'étude 2014, pour l'étude IFM 2009 cette proportion est de 74% avec l'étude IFM 2005 et de 77% avec IFM 2014).

D'autre part, et là est un point important des résultats et conclusions de cette étude, la proportion des patients recrutés par ces centres partagés (connue seulement pour les études finies que sont IFM 2005 et IFM 2009) et encore plus grande (ainsi 89% des patients recrutés lors de l'étude 2005 l'ont été par des centres partagés avec IFM 2009 et de 81% inversement). Les données de l'étude IFM 2014 pourraient également être utilisées, seulement l'objectif de l'étude entreprise ici est comme déjà mentionné de tester l'hypothèse de pertinence d'une étude de faisabilité dans un cadre bien défini, il est donc tout naturel qu'elles ne soient pas incorporer à l'analyse effectuée (de plus l'étude de faisabilité et sa modélisation dans le cas précis des hypothèses de proportions discutées ci-dessus fût mise en place au tout début de l'étude IFM 2014).

3.1.2 Modélisation des paramètres de la phase de recrutement

Une fois l'hypothèse des proportions des centres communs vérifiés et en supposant que la proportion de patients recrutés par ceux ci sera également grande, reste encore à définir la méthode d'estimation des paramètres qui vont servir pour modéliser la dynamique de recrutement avant qu'elle ne commence. Plusieurs approches sont alors étudiées par la suite.

Approche empirique

En considérant le i -ème centre de l'étude historique et en observant que n_i^h patients ont été recrutés depuis la date d'ouverture u_i du centre considéré.

Soit $\tau_i = T^h - u_i$ la durée totale d'activité de ce centre jusqu'à la fin de la phase de recrutement, alors l'estimateur du maximum de vraisemblance de l'intensité du centre i est :

$$\hat{\lambda}_i^h = \frac{n_i^h}{\tau_i^h}.$$

Approche bayésienne avec hypothèse gamma

Dès lors que les centres recrutent selon des intensités différentes entre eux, un modèle Poisson-gamma [10, 6] peut-être appliqué pour quantifier et incorporer au modèle cette variance. Les intensités λ_i^h sont toujours vu comme des réalisations d'une variable aléatoires distribuées selon une loi $\Gamma(\alpha, \alpha/\mu)$. Comme effectué en section 1.3.2, ces paramètres sont estimés par la méthode bayésienne empirique en utilisant la fonction log-vraisemblance du Théorème 1.

3.1.3 Comparaison de proximité des intensités entre centres partagés pour les deux approches.

La distinction est faite entre centres partagés et non partagés entre chaque étude, sont alors définis les ensembles de ces centres comme suit :

- **Centres partagés.** Soit $\mathcal{C}_s^h$ (respectivement $\mathcal{C}_s^n$) l'ensemble des centres de l'étude historique (respectivement la nouvelle étude) partagés entre les essais. Pour finir C_s^h (respectivement C_s^n) dénote son cardinal (forcément $C_s^h = C_s^n$).
- **Centres non partagés.** Soient $\mathcal{C}_u^n$ l'ensemble des centres qui sont impliqués dans le nouvel essai sans l'être dans l'essai historique et C_u^n son cardinal. Finalement soient $\mathcal{C}_u^h$ l'ensemble des centres qui sont impliqués dans l'essai historique mais pas dans le nouvel essai et C_u^h son cardinal.

Par la suite les paramètres relatifs aux centres partagés seront indexés par un s (respectivement un u pour les centres non partagés et un g pour l'ensemble des centres impliqués dans l'étude). De plus les index des sommes seront écrit selon les définitions d'ensembles précédentes. Les paramètres relatifs à l'essai historique auront pour leur part un exposant h (respectivement un exposant n pour les paramètres du nouvel essai).

Valeurs obtenues

Les précédentes définitions ont été appliquées aux données récoltées concernant les études IFM 2005 et IFM 2009, ces deux études étant terminées toutes les données de recrutement, dates d'ouvertures et durée réelle de recrutement sont disponibles.

Le tableau 3.2 montre selon les deux approches les valeurs obtenues, avec également des outils statistiques de comparaison tels les quantiles empiriques à 84 % pour l'approche empirique, les écarts-type pour l'approche bayésienne. Le paramètre μ représentant la moyenne de la loi gamma dont il découle permet une comparaison rapide avec les moyennes obtenues avec l'approche empirique.

Approche empirique

Une régression linéaire effectuée entre les différentes intensités pour les centres partagés entre IFM 2005 et IFM 2009 (analyse effectuée sur les paires $\{(\hat{\lambda}_i^h, \hat{\lambda}_{j(i)}^n), i \in \mathcal{C}_s^h\}$ où $j(i) \in \mathcal{C}_s^n$ est l'indice du centre i dans le nouvel essai) permet de comparer ces intensités en terme de proximité. En effet une fois la droite de régression obtenue il est possible de tester s'il existe une différence significative entre cette droite et la droite $y = x$ (intensités constantes pour toutes les études).

La droite de régression obtenue par l'explication des intensités de l'étude IFM 2009 par celles de IFM 2005 est :

$$\lambda_i^n = 0.44\lambda_i^h + 0.01$$

La corrélation est tout de même assez importante entre les intensités des deux études, et la positivité de la pente directrice permet de conclure que les centres qui recrutent beaucoup lors du premier essai vont recruter beaucoup aussi dans le deuxième essai (sauf quelques points visibles sur la Figure 3.1). Ainsi si une différence significative est trouvée entre la pente de la droite de régression et la pente de la droite $y = x$, alors une des études possédera forcément une moyenne d'intensités plus grande par rapport à l'autre étude, et donc les intensités ne pourront pas être proches d'une étude à l'autre.

TABLEAU 3.2 – Pour chaque catégorie de centre sont donnés la moyenne des intensités historiques, les quantiles empiriques à 84 % (équivalent à un écart-type) des intensités historiques, les valeurs (moyenne $\pm$ écart-type) des paramètres des loi gamma pour un essai en utilisant les données de l'autre essai (considéré comme historique). La seconde colonne considère IFM 2005 comme étude historique et IFM 2009 comme le nouvel essai, la troisième colonne considère la situation inversée. La dernière colonne représente les valeurs des p-values obtenues en effectuant des test de Wald entre les deuxième et troisième colonnes.

Paramètres	données historiques nouvel essai	IFM 2005 IFM 2009	IIFM 2009 IFM 2005	test de Wald
	Nombre de centres (nouveau)	67	76	
	Nombre de centres (hist.)	76	67	
	Recrutés (hist.)	611	693	
Globaux	$\frac{1}{C^h} \sum_{i=1}^{C^h} \hat{\lambda}_i^h$	0.012	0.017	0.62
	quantiles	[0.002 ; 0.016]	[0.003 ; 0.033]	
	α_G	1.75 ± 0.71	1.99 ± 0.84	
	μ_G	0.012 ± 0.027	0.017 ± 0.035	
	Nombre de centres (hist.)	50	50	
	Recrutés (hist.)	543	562	
Partagés	$\frac{1}{C^{h,s}} \sum_{i=1}^{C^{h,s}} \hat{\lambda}_i^h$	0.015	0.018	0.77
	quantiles	[0.003 ; 0.031]	[0.003 ; 0.034]	
	α_S	2.15 ± 1.02	2.15 ± 1.03	
	μ_S	0.015 ± 0.038	0.017 ± 0.042	
	Nombre de centres (hist.)	26	17	
	Recrutés (hist.)	68	131	
Non partagés	$\frac{1}{C^{h,u}} \sum_{i=1}^{C^{h,u}} \hat{\lambda}_i^h$	0.005	0.013	0.46
	quantiles	[0.002 ; 0.009]	[0.002 ; 0.023]	

En effet un test de comparaison de Student effectué entre la pente de la droite de régression avec 1 (la pente de la droite $y = x$) conclue à une différence significative entre ces deux valeurs (p-value de $1.7 \cdot 10^{-7}$). Ainsi de par l'homogénéité des proportions de recrutement (les gros centres recruteurs vont recrutés beaucoup et inversement), les intensités ne sont donc pas proches (dans un sens d'égalité stricte) d'une étude à l'autre (il faudrait des valeurs inférieurs et des valeurs supérieurs en même 'quantité' pour que le phénomène d'égalité se corrige).

Par conséquent l'idée naïve de reconduire les taux des centres partagés d'un essai à l'autre ne semble pas une bonne idée, hypothèse consolidée par les résultats obtenus présentés par la suite (proximité de la prédiction à la valeur réelle de la durée de la phase de recrutement).

La figure 3.1 représente pour chaque centre partagé ses deux intensités pour l'étude IFM 2005 (en abscisse) et IFM 2009 (en ordonnée), à quoi est rajouté la fonction de régression obtenue précédemment. Une grande dispersion autour de la droite de régression est visible, variance qui aura dans la majorité des cas pour conséquence l'ajout d'un biais sur la prédiction effectuée par la suite à partir de ces valeurs considérés comme paramètres incorporés au modèle prédictif.

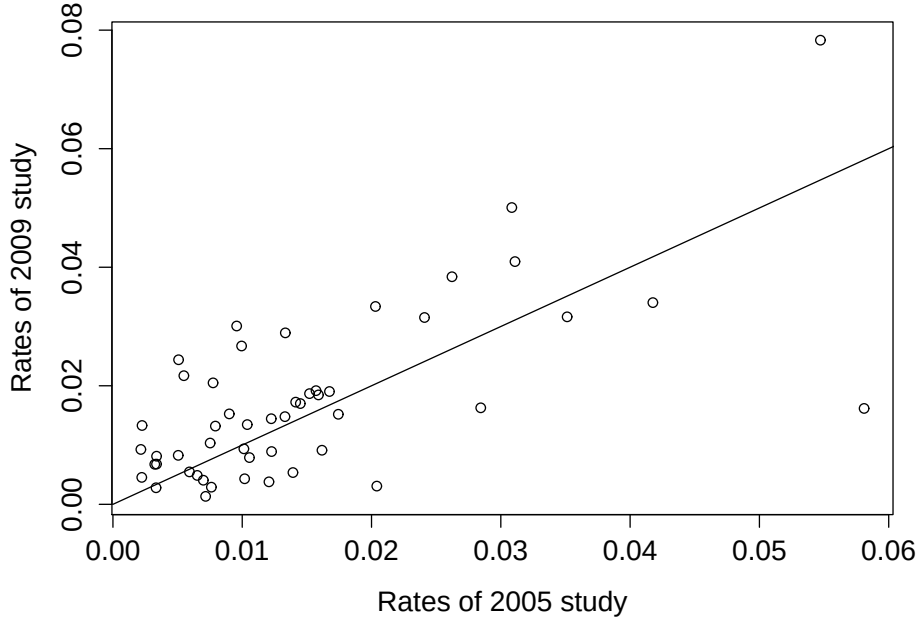


FIGURE 3.1 – Régression linéaire des intensités de recrutement des centres partagés entre les études IFM 2005 (en abscisse) et IFM 2009 (en ordonnée)

Approche bayésienne avec hypothèse gamma

En revanche, les comparaisons des distributions des paramètres (α_s^h, μ_s^h) et (α_s^n, μ_s^n) (calculées selon la méthode présentée en section 1.3.3 grâce à la matrice d'information global de Fisher) par des tests de Wald (grâce à la normalité asymptotique de ces distributions) ne montre pas de différences significatives. Ces comparaisons sont visibles sur la figure 3.2 et dans la table 3.2.

La figure 3.2 montre bien la proximité des régions de confiance entre les paramètres de IFM 2009 calculés à partir des données de IFM 2005 et inversement. Cette proximité est d'autant plus proche lorsque sont estimés les paramètres à partir des données des centres communs seulement (sous-figures b.).

Les paramètres relatifs aux centres non partagés ne sont pas ajoutés au graphique car l'estimation de ces paramètres $(\alpha_u^{h,n}, \mu_u^{h,n})$ selon les ensembles $\mathcal{C}_u^{h,n}$ n'est pas forcément stable selon les cardinaux C_u^h et C_u^n (s'ils sont inférieurs à 20). De plus ces centres sont souvent ceux dont les intensités sont les plus faibles, réduisant d'autant plus la pertinence des estimations (phénomène étudié au paragraphe 2.1.2).

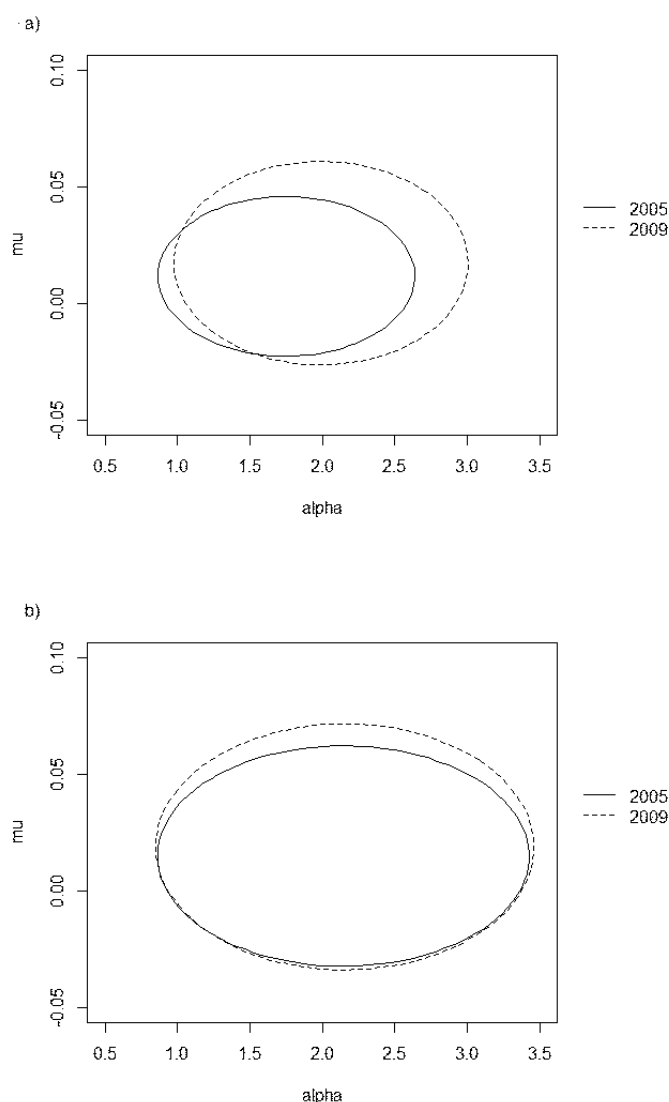


FIGURE 3.2 – Régions de confiance à 95% des paramètres (α, μ) . Sub-figure a) paramètres globaux . Sub-figure b) paramètres partagés.

L'idée de reconduire les distributions des vecteurs des intensités semble donc être une idée raisonnable, surtout pour la prédiction des dynamiques de recrutement des centres partagés. Seulement si pour les centres partagés les choses sont relativement claires, elles le sont beaucoup moins pour les nouveaux centres (réelle difficulté de l'étude de faisabilité effectuée ici, difficulté dont l'ampleur est justement réduite par la vérification des 'grandes' proportions entre facteurs partagées observées précédemment). En effet comme vu juste avant l'estimation des paramètres peut être très instables selon que la proportion des centres non partagés, mais le choix de la méthode à mettre en place pour considérer ces centres non partagés est très sensible aux études elles mêmes, et aucune forte corrélation n'est facilement observable et modélisable entre ces centres du fait de leur différences intrinsèques.

3.1.4 Prediction de la durée du recrutement pour le nouvel essai - u_i connus

Dans l'optique initiale d'apporter des recommandations sur la procédure à suivre pour mener des études de faisabilité dans le cadre précédemment défini, le choix des hypothèses de calcul des valeurs d'intensité est l'élément centrale de toute l'étude et le but est à partir d'elles de mesurer leur consistance en terme de pouvoir prédictif de la durée de recrutement sur des études finies et dont les dates d'ouvertures u_i sont toutes disponibles pour généraliser à des études dont le recrutement est encore à commencer.

Plusieurs scénarios ont ainsi été investigués et sachant que les études IFM 2005 et 2009 sont finies (les durées réelles sont donc connues et utilisables), les performances prédictives du modèle de faisabilité en terme de proximité de la prédiction par rapport à la durée réelle observée ont pu être comparées par des études comparatives entre elles deux : la durée de la phase de recrutement de l'étude IFM 2009 a été prédite à partir de l'estimation des paramètres elle même découlant des données historiques de l'étude IFM 2005 et inversement.

Scénarios de comparaison

Pour comparer les différentes hypothèses posées sur le choix des valeurs à assigner aux intensités de recrutement, ces scénarii se définissent comme suit :

- **Pour les centres partagés**, différents scénarios peuvent être explorés utilisant les informations de l'essai historique par :
 - **Hypothèse S1** : λ_i^n déterministe et égal à la valeur estimée de l'essai historique $\hat{\lambda}_i^h$,
 - **Hypothèse S2(g) (respectivement S2(s))** : λ_i^n distribué selon une loi $\Gamma(\alpha_g^h, \alpha_g^h/\mu_g^h)$ (respectivement $\Gamma(\alpha_s^h, \alpha_s^h/\mu_s^h)$).
- **Pour les centres non partagés**, le problème est plus compliqué. Aucune information n'est disponible directement sur ces centres et les hypothèses formulées doivent être faites sur leur potentiel de recrutement. Chaque hypothèse peut ainsi être liée à un modèle des intensités.
 - **Hypothèse U1(g) (respectivement U1(u))** : Aucune variabilité entre les centres n'est prise en compte. λ_i^n est choisi de manière déterministe et égale à la moyenne sur l'ensemble des valeurs des intensités de l'essai historique, $\frac{1}{C_g^h} \sum_{k \in C_g^h} \hat{\lambda}_k^h$ (respectivement égale à la moyenne des valeurs des intensités des centres non partagés de l'essai historique $\frac{1}{C_u^h} \sum_{k \in C_u^h} \hat{\lambda}_k^h$).
 - **Hypothèse U2** : Prise en compte de la variabilité entre les centres. λ_i^n est choisi aléatoirement selon une loi $\Gamma(\alpha_g^h, \alpha_g^h/\mu_g^h)$. L'utilisation des paramètres (α_u^h, μ_u^h) peut être pertinent cependant l'estimation est rarement précise du fait du peu de données disponibles.
 - **Hypothèse U3(L) (respectivement U3(M) et U3(H))** : Prise en compte de la variabilité entre les centres en considérant les λ_i distribués selon une loi gamma de paramètres (α_g^h, μ_g^h) restreint à une sous distribution $[0, q_1]$ (respectivement $[q_1, q_2]$ et $[q_2, +\infty]$) où q_1 et q_2 sont les 1/3 et 2/3 fractiles.

Ainsi les intensités sont restreintes aux basses (respectivement moyenne et haute) potentiel de recrutement.

Pour estimer la durée du nouvel essai, il est nécessaire d'évaluer l'intensité de recrutement globale :

$$t \rightarrow \Sigma^n(t) = \sum_{i \in \mathcal{C}_g^n} \lambda_i^n \max(t - u_i, 0),$$

où les intensités λ_i^n dépendent du choix du scénario considéré.

Le calcul des temps estimés pour les scénarios précédents se font de deux manières. La première méthode utilisée est celle de Monte-Carlo sur un nombre conséquent de générations (5000), la deuxième est une méthode analytique.

Méthode Monte-Carlo

Sachant un scénario choisi, la génération suivante est effectuée :

1. Génération des valeurs des intensités λ_i^n selon la distribution du scénario.
2. Génération (par le package NH poisson du logiciel R) de la trajectoire d'un processus non-homogène de Poisson X^n d'intensité Σ^n . Cette trajectoire est générée dans un intervalle $[0, 3 * T^n]$ assez grand pour couvrir avec une grande probabilité la durée de recrutement estimée.
3. Identification du temps de recrutement $\tilde{T}^n$ selon la trajectoire du processus X^n au premier instant tel que $X_{\tilde{T}}^n > N^n$.

5000 générations sont effectuées dont est tiré l'échantillon $\{\tilde{T}_j^n, j = 1, \dots, 5000\}$ permettant l'évaluation de la moyenne des durées de recrutement :

$$\hat{T}_{MC}^n = \frac{1}{5000} \sum_{j=1}^{5000},$$

avec l'intervalle de confiance à 95% $[\tilde{T}_{(125)}^n, \tilde{T}_{(4875)}^n]$ où l'échantillon $\{\tilde{T}_{(j)}^n, j = 1, \dots, 5000\}$ est l'échantillon ordonnée.

Méthode Analytique

Théorème 5 *Supposons que pour tout i , u_i est donné et notons $L^n = \max(u_i)$. Supposons aussi que $\mathbb{E}[\Sigma(L^n)] \ll N^n$.*

La durée de recrutement est alors estimée par :

$$\hat{T}_A^n = L^n + \frac{N^n - \mathbb{E}\Sigma(L^n)}{\sum_{i \in \mathcal{C}_g^n} \mathbb{E}\lambda_i^n}. \quad (3.1)$$

Pour un nombre C^n assez grand ($C^n > 10$), $[T_*(\delta, N^n), T^*(\delta, N^n)]$ peut servir d'intervalle de confiance prédictif à un niveau δ de T^n où $T_*(\delta, N^n)$ (resp. $T^*(\delta, N^n)$) est une solution selon la variable t des équations :

$$\mathbb{E}\Sigma^n(t) + z_{1-\delta/2}\sqrt{\mathbb{E}\Sigma^n(t) + \mathbb{V}\Sigma^n(t)} = N^n \quad \left(\text{resp. } \mathbb{E}\Sigma^n(t) - z_{1-\delta/2}\sqrt{\mathbb{E}\Sigma^n(t) + \mathbb{V}\Sigma^n(t)} = N^n \right).$$

Proof : Pour un processus de Cox d'intensité $\Sigma^n(t)$, pour tout $t > 0$,

$$\begin{aligned} \mathbb{E}X^n(t) &= \mathbb{E}\Sigma^n(t), \\ \mathbb{V}X^n(t) &= \mathbb{E}\Sigma^n(t) + \mathbb{V}\Sigma^n(t). \end{aligned}$$

De plus pour tout $t > L$, l'intensité cumulée est :

$$\Sigma^n(t) = \Sigma^n(L^n) + (t - L^n) \sum_{i \in \mathcal{C}_g^n} \lambda_i^n$$

Ainsi,

$$\mathbb{E}X(t) = \mathbb{E}\Sigma^n(L^n) + (t - L) \sum_{i \in \mathcal{C}_g^n} \mathbb{E}\lambda_i^n.$$

La moyenne de recrutement est approximativement le temps tel que $\mathbb{E}X(T^n) = N^n$, et ainsi il est approximé par (3.1). L'intervalle de confiance est calculé selon le Lemme 2.1 de [6] en notant qu'il n'y a pas de temps intermédiaire ici ($t_0 = 0$). $\square$

3.1.5 Résultats d'analyse

IFM 2009 estimé par les données d'IFM 2005

Les résultats des comparaisons des deux approches (empirique contre bayésienne) nous pousse à préférer l'approche bayésienne. Pour appuyer cette réflexion une comparaison en terme de proximité de la prédiction du temps de recrutement de chaque approche à la durée réelle est effectuée en considérant l'étude IFM 2009 comme la nouvelle étude et IFM 2005 comme l'étude historique servant au paramétrage du modèle. Les résultats obtenus par la méthode Monte-Carlo et la méthode analytique sont récoltés dans le tableau 3.3 et illustrés par les Figures 3.3-3.4.

Le tableau présente pour chaque hypothèse posée sur les centres non partagés (U1 U2 ou U3) les résultats pour chaque choix concernant les centres partagés (S1 ou S2). La durée moyenne ainsi que son intervalle de confiance à 95 % sont estimés, et sachant que la durée de recrutement de l'étude IFM 2009 est connue ($T^n = 746$ jours), la performance prédictive peut être étudiée grâce à l'erreur relative de prédiction définie comme :

$$\widehat{\text{RE}}_{\text{MC}}^n = \frac{|\widehat{T}_{\text{MC}}^n - T^n|}{T^n} \quad \text{resp.} \quad \widehat{\text{RE}}_{\text{A}}^n = \frac{|\widehat{T}_{\text{A}}^n - T^n|}{T^n}.$$

TABLEAU 3.3 – Moyennes des durée prédites pour l’essai IFM 2009 en utilisant IFM 2005 comme étude historique pour les différentes hypothèses concernant les centres paratgés et non paratgés. La durée réelle est de $T^n = 746$ jours.

Non partagés Hypothèse	Partagés Choix	méthode Monte Carlo			méthode Analytic		
		Prédictions	Ecart type	Erreur rel.	Prédictions	Ecart type	Erreur rel.
U1(g)	S1	915.62	29.81	0.23	881.24	28.31	0.18
	S2(g)	1066.11	81.87	0.43	1046.39	81.35	0.40
	S2(s)	911.08	63.67	0.22	891.70	64.86	0.20
U1(u)	S1	1019.02	33.27	0.37	980.33	32.07	0.31
	S2(g)	1226.97	112.25	0.64	1194.61	108.40	0.60
	S2(s)	1017.32	80.43	0.36	990.32	81.66	0.33
U2	S1	913.39	42.02	0.22	876.82	40.83	0.18
	S2(g)	1062.97	99.82	0.42	1039.96	92.05	0.39
	S2(s)	906.39	66.92	0.21	887.29	71.17	0.19
U3(L)	S1	1046.62	34.68	0.40	1001.70	33.77	0.34
	S2(g)	1266.75	119.22	0.70	1227.61	115.76	0.65
	S2(s)	1035.85	84.44	0.39	1011.54	85.95	0.36
U3(M)	S1	946.19	28.48	0.27	904.44	30.00	0.21
	S2(g)	1106.62	85.01	0.48	1080.40	87.81	0.44
	S2(s)	943.10	64.90	0.26	914.82	68.98	0.23
U3(H)	S1	785.55	23.45	0.05	763.80	30.61	0.02
	S2(g)	898.92	55.91	0.20	880.28	62.47	0.18
	S2(s)	774.92	45.75	0.04	774.25	51.71	0.04

Nous remarquons premièrement qu'importe les hypothèses considérées, la durée de recrutement est toujours sur-estimée. Ce phénomène n'est pas surprenant sachant que les études sont similaires et qu'une sélection de centres a été effectuée entre les études IFM 2005 et IFM 2009 pour garder les meilleures centres. Ainsi les centres de 2009 recrutent en moyenne plus (donc les prédictions à partir des paramètres estimés grâce aux données de l'étude IFM 2005 vont être plus grandes).

Il est intéressant de noter que pour toutes les hypothèses concernant les centres non partagés, le scénario avec l'erreur de prédiction la plus faible est celui utilisant les paramètres estimés sur les données des centres partagés (les scénarios S1 et S2(s) donnent des résultats équivalents, et tous deux meilleurs que S2(g)). Ces résultats ne sont pas surprenant en observant que les moyennes des intensités de recrutement de ces deux scénarios sont aussi équivalents. Cependant nous observons une différence dans la taille des intervalles de confiance à 95% entre ces deux scénarios. En effet le scénario S2(s) possède des intervalles de confiance plus grand, du fait de l'aléa venant des intensités qui sont variables contrairement à S1. Ceci pousse à préférer le scénario 2 comme souligné par Senn [56]. En effet avoir un intervalle de confiance le plus petit possible n'est pas un objectif ici. Le but est d'offrir une approximation ainsi qu'un intervalle de confiance qui puisse couvrir la zone dans laquelle se trouve la vraie valeur. Des petits intervalles de confiance limitent cette hypothèse et ne permettent pas de surmonter le possible biais. Le choix le plus réaliste concernant les centres partagés semble donc être le choix S2(s).

En excluant les hypothèses U3(L), U3(M) et U3(H) et en considérant tous les choix concernant les centres partagés, les meilleures hypothèses sur les centres non partagés sont les hypothèses U1(g) et U2. Leur durée moyenne sont les plus proches de la durée réelle, mais en considérant la remarque précédente sur la taille des intervalles de confiance, on va préférer l'hypothèse U2 qui fournit les intervalles de plus grande taille. L'hypothèse U1(u) donne les plus mauvais résultats, ce qui confirme l'idée qu'utiliser les données provenant des centres non partagés pour une étude où justement les centres partagés sont nombreux est une mauvaise idée.

En ce qui concerne les hypothèses U3, la meilleure hypothèse est l'hypothèse U3(H) en terme de proximité par rapport à la durée réelle. Ceci peut s'expliquer par le fait que considérer les centres à haut potentiels de recrutement permet de contrer le biais provenant de la sélection de centres. Finalement les résultats obtenus pour ces trois hypothèses souligne le lien entre l'hypothèse sur les intensités et la prédiction de la durée estimée.

IFM 2005 estimé par les données d'IFM 2009

L'étude peut être enrichie en considérant l'étude IFM 2009 comme l'étude historique et IFM 2005 comme la nouvelle étude. Les résultats sont présentés de la même manière que précédemment, dans le Tableau 3.4 sont collectés tous les résultats qui sont illustrés dans la Figure 3.5. La durée réelle de l'étude IFM 2005 est $T^n = 914$ jours.

Il est important de noter que le contexte est très différent de l'étude précédente. En effet cette fois les nouveaux centres sont ceux qui ont été enlevés entre IFM 2005 et IFM 2009, centres qui sont majoritairement des centres à faible potentiel de recrutement.

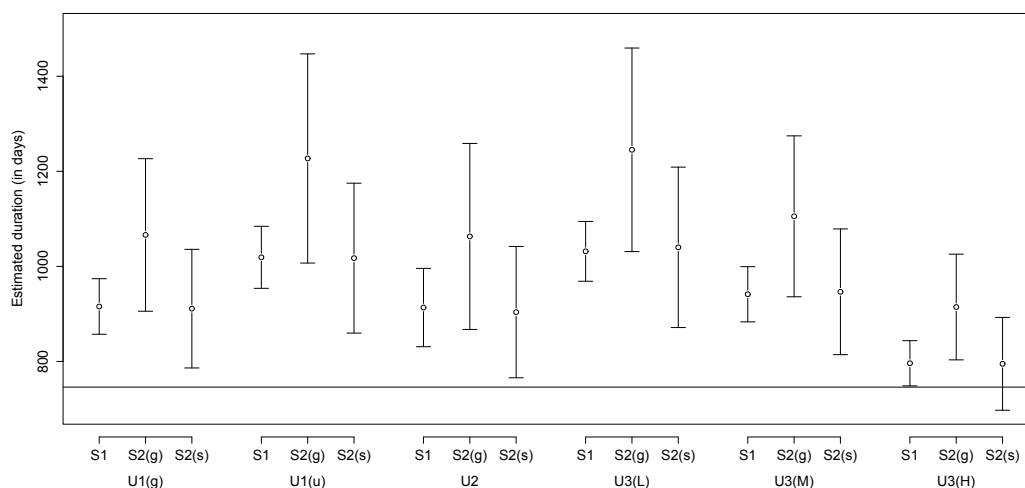


FIGURE 3.3 – Méthode de Monte-Carlo : estimations et intervalles de confiance à 95% des estimations de la durée de l'essai IFM 2009 grâce aux données de l'essai IFM 2005. Le trait plein désigne la vraie durée de l'essai IFM 2009.

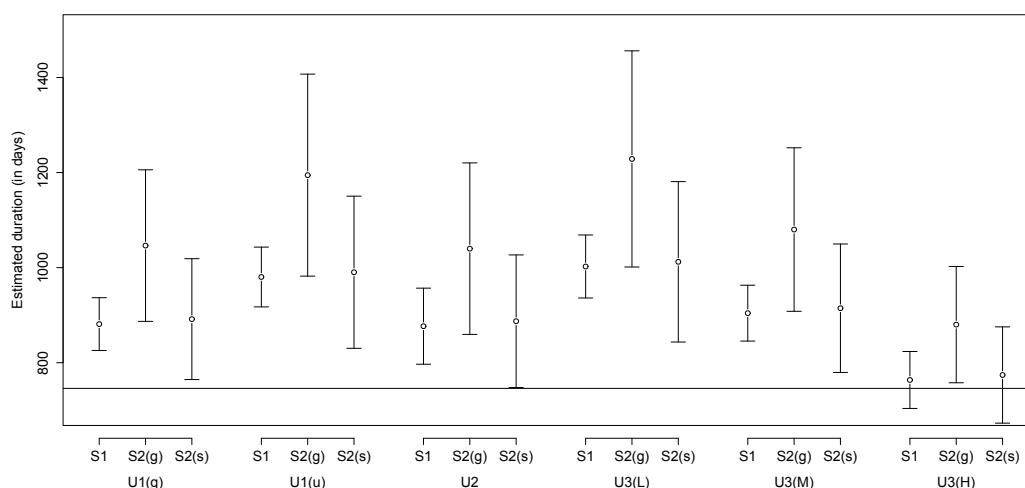


FIGURE 3.4 – Méthode analytique : estimations et intervalles de confiance à 95% des estimations de la durée de l'essai IFM 2009 grâce aux données de l'essai IFM 2005. Le trait plein désigne la vraie durée de l'essai IFM 2009.

De plus l'étude IFM 2009 visait à étudier en plus des effets d'un traitement les effets d'une intervention chirurgicale (intervention ne pouvant se faire que dans des centres possédant les installations nécessaires et donc souvent plus conséquent que des centres ne s'occupant que de l'étude du traitement). Le but de cette étude est donc d'étudier les hypothèses sur les centres non partagés plutôt que de la comparaison des choix sur les centres partagés. Ainsi seul les scénarios S1 et S2(s) sont considérés.

TABLEAU 3.4 – Moyennes des durée prédites pour l'essai IFM 2005 en utilisant IFM 2009 comme étude historique pour les différentes hypothèses concernant les centres paratgés et non paratgés. La durée réelle est de $T^n = 914$ jours.

Non partagés	Partagés	méthode Monte Carlo			méthode Analytic		
Hypothèse	Choix	Prédictions	Ecart type	Erreur rel.	Prédictions	Ecart type	Erreur rel.
U1(g)	S1	723.89	17.45	0.21	688.54	18.01	0.25
	S2(s)	718.92	39.97	0.21	709.68	39.99	0.22
U1(u)	S1	756.28	19.05	0.17	716.75	19.66	0.22
	S2(s)	756.59	47.86	0.17	739.44	45.55	0.19
U2	S1	723.85	25.01	0.21	687.06	23.76	0.25
	S2(s)	720.97	44.16	0.21	708.11	43.22	0.23
U3(L)	S1	813.07	23.14	0.11	769.10	23.28	0.16
	S2(s)	812.44	55.99	0.11	794.48	57.11	0.13
U3(M)	S1	742.43	19.05	0.19	703.18	19.24	0.23
	S2(s)	735.38	45.38	0.20	757.32	45.99	0.17
U3(H)	S1	651.34	15.28	0.29	623.90	16.8	0.32
	S2(s)	663.75	32.54	0.27	641.23	30.19	0.30

Premièrement il n'est pas surprenant de voir que la durée réelle est sous-estimée. C'est un contre-coup de la sélection de centres entre IFM 2005 et IFM 2009.

Les hypothèses sur les centres non partagés amènent à des résultats équivalents. Ce n'est pas surprenant sachant que les estimations des paramètres globaux, partagés et non partagés sont proches (voir Tableau 3.2), donc qu'importe l'hypothèse considérée, les résultats seront proches. Les hypothèses provenant de l'approche bayésienne proposent cependant des intervalles de confiance plus réalistes.

Enfin on observe que entre les hypothèses U3, l'hypothèse U3(L) possède les meilleurs résultats, ceci corrobore l'idée que les centres à petit potentiel de recrutement ont été écartés entre IFM 2005 et IFM 2009.

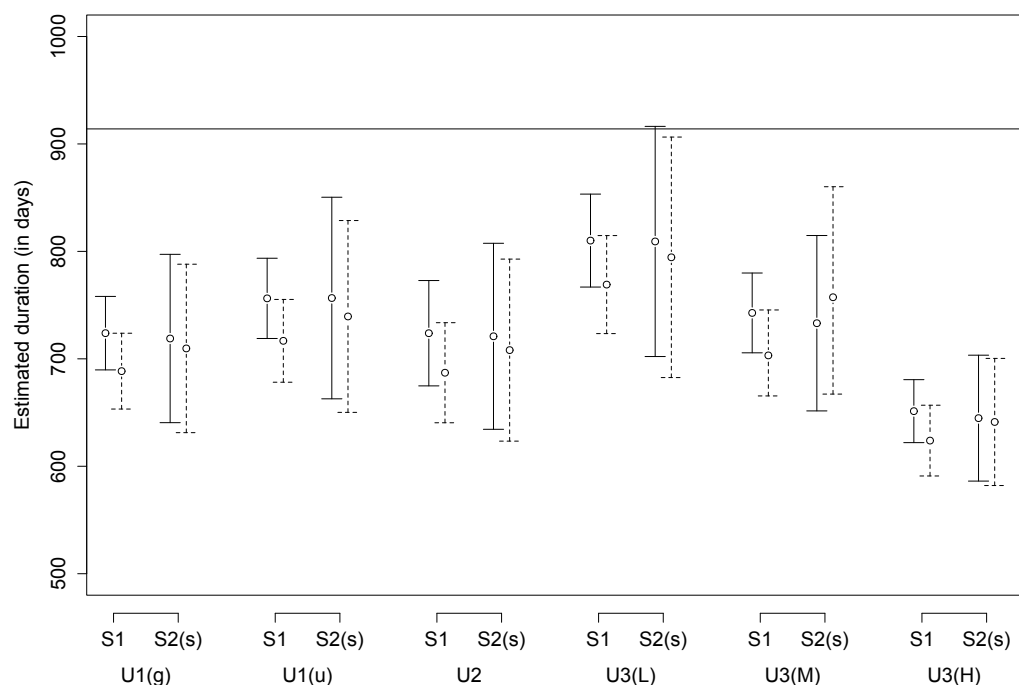


FIGURE 3.5 – Méthode de Monte-Carlo : estimations et intervalles de confiance à 95% des estimations de la durée de l'essai IFM 2005 grâce aux données de l'essai IFM 2009. Le trait plein désigne la vraie durée de l'essai IFM 2005.

3.1.6 Recommandations et exemple illustratif - Faisabilité de l'essai IFM 2014

Recommandations

Nous donnons ici des recommandations selon les résultats des paragraphes précédents. Premièrement la méthodologie présentée ici n'est valable que dans le cas où l'essai fini est similaire à l'essai en cours ou à venir que ce soit en terme de domaine thérapeutique, les critères d'inclusion/exclusion, régions et population cible. Les résultats poussent à préférer l'utilisation du modèle Poisson-gamma pour évaluer la faisabilité d'un essai clinique dans les conditions précitées.

La modélisation des centres partagés et des centres non partagés ne sera pas la même et dépendra de l'étude considérée.

Un point important est que la pertinence de la méthode augmente avec la proportion de centres partagés.

- Pour les centres partagés, il est suggéré de modéliser les intensités de recrutement par le choix $S2(s)$.
- Pour les centres non partagés, il est suggéré d'utiliser l'hypothèse gamma. Le choix entre les hypothèses $U2$, $U3(L)$, $U3(M)$, et $U3(H)$ dépend de l'étude et de possibles paramètres donnés à priori par un expert.

Finalement la modélisation des dates d'ouverture des centres est à considérée. Les résultats de [3] et [46] encourage l'utilisation d'une stratégie uniforme.

Exemple illustratif - faisabilité de l'essai IFM 2014

La méthodologie présentée dans ce paragraphe est appliquée pour estimer la durée du recrutement de l'essai IFM 2014. La durée attendue par les investigateurs est $T^n = 730$ jours. Le contexte de cette étude est résumé dans le Tableau 3.1. Les données historiques utilisées seront celles provenant de l'essai IFM 2009 (plus proche en terme de proportion de centres partagés et de domaine thérapeutique).

Le Tableau 3.5 présente les estimations des paramètres qui seront utilisés pour la prédiction.

L'étude IFM 2014 est toujours en cours de réalisation. Les dates d'ouvertures ne sont pas utilisées pour la prédiction du temps de recrutement de l'essai. La stratégie uniforme est utilisée pour modéliser ces dates d'ouvertures. Les résultats des prédictions sont collectés dans le Tableau 3.6 et illustrés par la Figure 3.6.

Comme pour les études précédentes, les hypothèses $U1(g)$ et $U2$ donnent des résultats équivalents que ce soit en terme de moyenne prédites qu'en terme d'écart type. Ce dernier point est dû à la faible proportion des centres non partagés qui réduit l'ajout d'aléa de l'hypothèse $U2$ par rapport à l'hypothèse $U1(g)$. Finalement, il est intéressant de noter que les résultats des hypothèses $U2$, $U3(L)$ et $U3(M)$ sont très proches tandis que ceux de l'hypothèse $U3(H)$ sont plus petits, ceci montre l'importance des centres à grands potentiels de recrutement (car malgré la petite proportion de centres non partagés cette hypothèse arrive à donner des résultats assez éloignés).

Une discussion avec l'équipe s'occupant de l'essai nous pousse à considérer l'hypothèse $U3(M)$ pour les centres non partagés (centres avec un potentiel moyen de recrutement). Les résultats ainsi obtenus donnent une valeur prédites de 577 jours et un intervalle de confiance à 95% de [505, 650] jours.

A l'instant de rédaction de cette partie, 20 patients sont encore à recruter, les 280 patients déjà recrutés l'ont été en 560 jours. Notre stratégie semble donc appropriée.

TABLEAU 3.5 – Pour chaque catégorie de centre sont donnés la moyenne des intensités historiques, les quantiles empiriques à 84 % (équivalent à un écart-type) des intensités historiques, les valeurs (moyenne $\pm$ écart-type) des paramètres des loi gamma pour un essai en utilisant les données IFM 2009 comme données historique.

Paramètres	données historiques	IFM 2009
	nouvel essai	IFM 2014
Globaux	Nombre de centres (nouveau)	47
	Nombre de centres (hist.)	67
	Recrutés (hist.)	693
	$\frac{1}{C^h} \sum_{i=1}^{C^h} \hat{\lambda}_i^h$	0.016
	quantiles	[0.003 ; 0.033]
	α_G	1.99 ± 0.84
	μ_G	0.017 ± 0.035
Partagés	Nombre de centres (hist.)	41
	Recrutés (hist.)	534
	$\frac{1}{C^{h,s}} \sum_{i=1}^{C^{h,s}} \hat{\lambda}_i^h$	0.021
	quantiles	[0.002 ; 0.038]
	α_S	2.31 ± 1.25
	μ_S	0.022 ± 0.050
Non partagés	Nombre de centres (hist.)	6
	Recrutés (hist.)	159
	$\frac{1}{C^{h,u}} \sum_{i=1}^{C^{h,u}} \hat{\lambda}_i^h$	0.011
	quantiles	[0.003 ; 0.020]
	α_U	3.45 ± 1.96
	μ_U	0.014 ± 0.047
1/3-fractiles	q_1	0.0103
	q_2	0.0198

3.1.7 Conclusion

Dans le cadre d'une étude de faisabilité d'un essai clinique, ce paragraphe préconise l'emploi du modèle Poisson-gamma pour estimer la durée du recrutement de l'essai en se basant sur les paramètres estimés à partir de données d'un essai cliniques déjà fini. En effet cette approche offre des résultats plus convaincant que la méthode d'attribution déterministe des intensités historiques de recrutement. La prédiction du temps de recrutement par le modèle Poisson-gamma est au moins aussi bonne que la méthode déterministe mais bénéficie d'intervalles de confiance plus réalistes.

TABLEAU 3.6 – Moyenne des durées prédites pour l’essai clinique IFM 2014 avec IFM 2009 comme essai historique pour différentes hypothèses concernant les centres non partagés. Choix des centres partagés S2(s).

Hypothèse des centres non partagés	Durée moyenne	Ecart type
U1(g)	586.36	44.33
U1(u)	662.45	44.05
U2	583.43	45.30
U3(L)	584.91	35.82
U3(M)	577.32	36.83
U3(H)	534.12	33.97

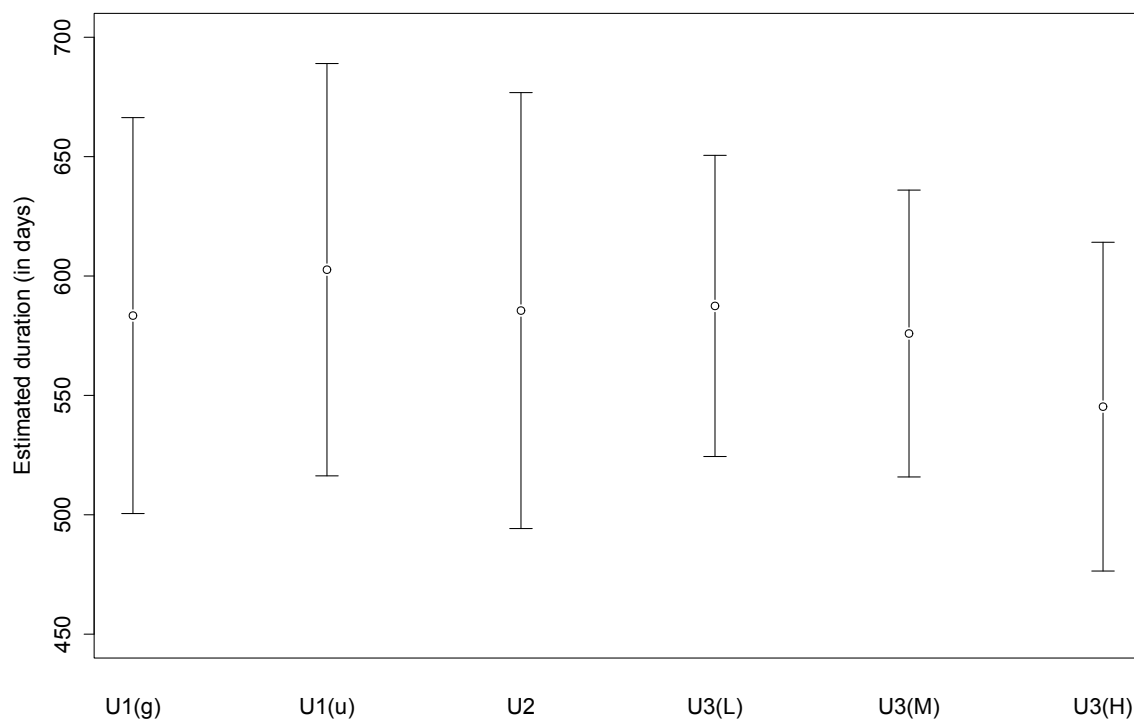


FIGURE 3.6 – Intervalles de confiance à 95% et durée moyenne des prédictions pour les différentes hypothèses considérées de l’essai IFM 2014 avec IFM 2009 comme étude historique.

Des méthodes de prédictions ont été présentées avec le calcul du temps de recrutement par les méthode Monte-Carlo et analytique offrant des résultats pertinents et similaires.

Il est également important de noter que les performances du modèle Poisson-gamma augmentent avec la proportion de centres partagés. La qualité d’estimation des paramètres globaux et partagés augmente en effet avec la proportion de centres partagés. C’est un point très important car en pratique les centres qui sont réassignés sont ceux présentant des bons comportement de recrutement.

Les recommandations présentées dans cette partie propose une méthodologie pertinente pour l'utilisation de données historiques dans le cas d'essais cliniques proches en terme de domaine thérapeutiques, de critères d'inclusion/exclusion, régions et population cible.

3.2 A la prédiction dynamique selon des critères de jugements sur le nombre de patients

3.2.1 Randomisation : impact sur l'essai

La randomisation est essentielle lors de la planification d'une essai clinique. Les propriétés de multiples méthodes de randomisation ont été étudiées depuis maintenant 30 ans. Nous pouvons citer entre autres les travaux de Hallstrom et Davis [31], Matts et Lachin [43], Lachin [40, 39], Pocock [52], Rosenberger et Lachin [53], Senn[56] et Hedden et al. [32].

Selon la méthode de randomisation utilisée, l'impact sur la puissance du test statistique ainsi que sur l'approvisionnement des traitements est à considérer. En effet les patients seront différemment subdivisés en plusieurs sous-groupes selon l'analyse menée, sous-groupes pouvant entre autre être source de biais de déséquilibre (ce cas est par exemple étudié par Hallstrom et Davis dans [31] dans le cas d'une randomisation stratifiée par bloc de permutation). Il est également possible d'observer selon la méthode de randomisation considérée une baisse de la puissance du test statistique effectué (Lachin l'étudie dans [40] dans le cas d'une randomisation par bloc de permutations).

Les méthodes considérées dans ce paragraphe seront celles envisagées par Anisimov dans [4], à savoir la méthode de randomisation non-stratifiée (la plus simple) et la méthode de randomisation centre-stratifiée (l'équilibrage porte sur les effectifs des centres ou région de centres) par bloc de permutation. Ces deux méthodes sont en générale celles utilisées en recherche clinique. Les méthodes avec stratifications permettent certes de contrer l'effet de covariables sur la puissance du test statistique, mais dans le cas d'essai clinique dont l'effectif est de taille moyenne voire large [40, 32], cet impact est négligeable comparativement à une méthode non-stratifiée. Dans l'optique d'étudier les impacts de ces deux méthodes selon plusieurs critères de jugement envisagés par la suite, un rappel de ces deux méthodes est donné.

Nous nous plaçons dans le cas d'un essai dont l'objectif de recrutement est de n patients sollicitant C centres, et potentiellement s régions comportant chacune C_s centres. Considérons que l'essai implique K traitements avec comme vecteur d'allocation $(k_1, \dots, k_K)$ et $K_1 = \sum_{i=1}^K k_i$ la taille du bloc. Soient $n_{i,j}$ la prédiction du nombre de patients randomisés dans le centre i pour le traitement j (resp. $n_j(I_s)$ la prédiction du nombre de patients randomisés dans la région I_s pour le traitement j). Plusieurs assumptions sont faites pour simplifier l'étude, à savoir la date d'ouvertures des centres est communes et égales à 0, et que $M = n/K_1$ est entier.

Randomisation non-stratifiée

La randomisation non stratifiée (pour l'effet des centres) consiste en l'affectation aléatoire des patients par bloc sans considération du centre d'origine du patient. Cette méthode permet de minimiser le déséquilibre entre les groupes pour l'essai global, mais augmente le déséquilibre par région ou par centre en comparaison de la méthode de stratification par centre/région.

Comme montré dans [34], la distribution des patients d'une région ou d'un centre peut être vu comme beta-binomiale. par la suite nous noterons par $P(n, C, C_s, \alpha, k)$ pour tout $k = 0, \dots, n$ la probabilité :

$$\mathbb{P}[n(I_s) = k] = \binom{n}{k} \frac{\mathcal{B}(\alpha C_s + k, \alpha(C - C_s) + n - k)}{\mathcal{B}(\alpha C_s, \alpha(C - C_s))}, \quad (3.2)$$

où $\mathcal{B}(a, b) = \int_0^1 x^{a-1} (1-x)^{b-1} dx$ et,

$$\mathbb{E}[n(I_s)] = nC_s/C, \quad \mathbb{V}[n(I_s)] = \frac{nC_s(C - C_s)(\alpha C + n)}{C^2(\alpha C + 1)}$$

Pour un traitement j nous avons Mk_j patients au total dans ce groupe, ainsi pour une région I_s :

$$\mathbb{P}[n_j(I_s) = i_j, j = 1, \dots, K] = \prod_{j=1}^K P(Mk_j, C, C_s, \alpha, i_j) \quad (3.3)$$

Les équations (3.2) et (3.3) permettent pour tout niveau de confiance $1 - \epsilon$ de calculer la limite supérieure de prédiction L_j pour le nombre de patients dans la région I_s pour tous les traitements de sorte que :

$$\mathbb{P}[n_j(I_s) \leq L_j, j = 1, \dots, K] \geq 1 - \epsilon$$

Stratification par centre/région

Dans cette méthode la randomisation se fait par centre et non plus de manière globale. Cette méthode possède l'avantage de réduire le déséquilibre par centre/région, mais augmente ce déséquilibre lorsque l'on considère l'échantillon global. Pour chaque centre une liste différente de randomisation est générée selon une randomisation par bloc de permutation. Il est ainsi possible pour un centre ou une région de posséder un nombre de patients dont n n'est pas un multiple, laissant certains bloc incomplets, pouvant causer un déséquilibre dans ce centre du nombre de patients par traitement. Ainsi dans le cas d'un essai multicentrique, de multiples blocs incomplets peuvent causer un déséquilibre du nombre total de patients par groupe de traitements et donc une baisse de la puissance du test.

Soit un bloc incomplet de taille m ($m = 1, \dots, K_1 - 1$) dans le centre i . Notons $\xi_j(m)$ le nombre d'occurrences du traitement j dans ce bloc incomplet. Alors $\xi_j(m)$ possède

une distribution hypergéométrique :

$$\mathbb{P}[\xi_j(m) = l] = \binom{k_j}{l} \binom{K_1 - k_j}{m - l} \binom{K_1}{m}^{-1}, l = 0, 1, \dots, \min(k_j, m).$$

Dans le cas où $m = 0$, nous posons $\xi_j(0) = 0$. Soit $\text{int}(a)$ la partie entière de $a \in \mathbb{R}$ et $\text{mod}(a, k) = a - \text{int}(a/k)k$, alors,

$$n_{i,j} = \text{int}(n_i/K_1)k_j + \xi_j(\text{mod}(n_i, K_1)).$$

Comme montré dans [4], il est alors possible dans le cas d'un recrutement dont tous les centres sont initialisés au début de l'étude d'approcher la loi de $n_j(I_s)$ par :

$$n_j(I_s) = C_s A(j) + \sqrt{C_s} B(j) \mathcal{N}(0, 1),$$

où $\mathcal{N}(0, 1)$ est une variable normale centrée réduite, $A(j) = \mathbb{E}[n_{i,j}]$ et $B^2(j) = \mathbb{V}[n_{i,j}]$.

Impact sur la puissance de test et de la taille de l'échantillon

Dans le but d'obtenir une étude fiable, une répartition identique entre les différents groupes randomisés est préférable d'un point de vue statistique. Il semble ainsi préférable de privilégier la méthode non stratifiée qui assure une répartition égale entre chaque groupe. Cependant Anisimov [4] montre que dans le cas d'un essai multicentrique, l'impact statistique d'une mauvaise répartition est assez faible comparé à la taille de l'échantillon total. De plus pour des test simples de comparaison de moyennes, l'apport d'une augmentation de la taille de l'échantillon du a un déséquilibre lors de la randomisation est statistiquement négligeable en ce qui concerne la puissance du test mené.

Ainsi l'impact de la randomisation dans le cas où l'effet centre n'est pas observable n'affecte pas les propriétés de l'analyse statistique menée, comme nous allons le voir par la suite, mais a un impact non négligeable sur la chaîne d'approvisionnement des traitements.

Ainsi nous nous plaçons dans le cas d'une randomisation stratifiée par centre, et nous allons étudier les impacts de cette méthode de randomisation sur la puissance du test statistique ainsi que sur la taille de l'échantillon de l'essai.

Plaçons nous dans le cas simple où seuls deux traitements a et b sont comparés. Notons n_a (resp. n_b) le nombre total de patients randomisés dans le groupe du traitement a (resp. du traitement b), et $\Delta = n_a - n_b$ le déséquilibre total du nombre de patients entre les deux traitements.

Pour évaluer l'augmentation nécessaire de la taille de l'échantillon en vue de compenser la perte de puissance du au déséquilibre de randomisation, nous nous plaçons dans le cas simple d'un test de comparaison de moyennes entre 2 groupes. Supposons que n patients recrutés par C centres sont randomisés dans les deux groupes avec comme taille de bloc K_1 . Notons $\{x_1, \dots, x_{n_a}\}$ pour le traitement a et $\{y_1, \dots, y_{n_b}\}$ pour le traitement b les réponses au traitement des patients de chaque groupe. Supposons que les

observations sont indépendantes, de moyennes m_a et m_b inconnues et de variance σ^2 connue.

Dans le cas d'une étude sans effet de centre, nous allons voir que la randomisation par stratification selon les centres ne change pas significativement les résultats de l'analyse par rapport à une randomisation non stratifiée.

En effet sachant que pour tester l'hypothèse $\mathcal{H}_0 : m_a - m_b = 0$ contre $\mathcal{H}_1 : m_a - m_b \geq h$ avec probabilité δ et γ d'erreur de première et second espèce, les valeurs n_a et n_b doivent satisfaire l'équation :

$$\frac{h}{\sigma \sqrt{\frac{1}{n_a} + \frac{1}{n_b}}} = z_{1-\delta/2} + z_{1-\gamma}$$

Dans le cas d'une randomisation équilibré où $n_a = n_b = n/2$ (en supposant n pair), la taille de l'échantillon nécessaire n_{bal} est donc $n_{bal} = 4\sigma^2(z_{1-\delta/2} + z_{1-\gamma})^2/h^2$.

Dans le cas où Δ n'est pas nul, la taille de l'échantillon doit être ajustée pour avoir la même puissance de test. Comme montré dans [4], l'augmentation nécessaire de la taille de l'échantillon est égale à $n_+ = s^2(n_{bal}, C, K_1, \alpha)C/n_{bal}$, où $s^2(n_{bal}, C, K_1, \alpha)$ est donné par :

$$s^2(n_{bal}, C, K_1, \alpha) = \frac{1}{K_1 - 1} \sum_{m=1}^{K_1-1} m(K_1 - m)q_m(n, C, \alpha, K_1),$$

$$q_m(n, C, \alpha, K_1) = \sum_{l=0}^{C/K_1-1} P(n, C, 1, \alpha, m + lk_1).$$

Pour des valeurs de C et n assez grande, nous pouvons utiliser l'approximation $s^2 = (K_1 + 1)/6$, sachant de plus que $C \leq n_{bal}$ et que pour deux traitements $K_1 = 4$, le résultat précédent implique que $n_+ \leq 2$. Ainsi dans le cas d'une comparaison de deux traitements et en supposant qu'aucun autre facteur de stratification n'agit, les deux méthodes de stratifications sont pratiquement équivalentes.

3.2.2 Indicateurs de performances de la dynamique de recrutement

La prédiction faite par le modèle Poisson-gamma de la durée nécessaire au recrutement des n patients par les C centres sollicités se base sur l'estimation des paramètres de la dynamique globale de recrutement grâce à la propriété d'additivité des processus singulier de chaque centre. Cependant certains paramètres ou variables rentrant dans le modèle peuvent être utilisés pour étudier d'autres aspects du recrutement. Nous avons ainsi vu qu'une étude de sensibilité peut-être menée sur les paramètres estimés par le modèle 1.3.3 ou de connaître la probabilité de finir à temps le recrutement (cf Théorème 3). L'objet de ce paragraphe est de présenter des études se basant sur le calcul d'indicateurs de performances de la dynamique de recrutement à partir du modèle Poisson-gamma selon les objectifs définis lors de la mise en place de l'essai.

Cartes de contrôle du recrutement

Le modèle Poisson-gamma permet une bonne prédiction de la durée de recrutement d'un essai clinique. L'autre avantage de cette modélisation est qu'il offre la possibilité de calculer des intervalles de confiance et de mettre en place des cartes de contrôle permettant un suivi dynamique du recrutement. Un exemple en est donné à la figure 3.7 où deux cartes sont créées à partir des données recueillies en deux instant d'analyse $t_1 = 1$ an et $t_1 = 1.5$ ans.

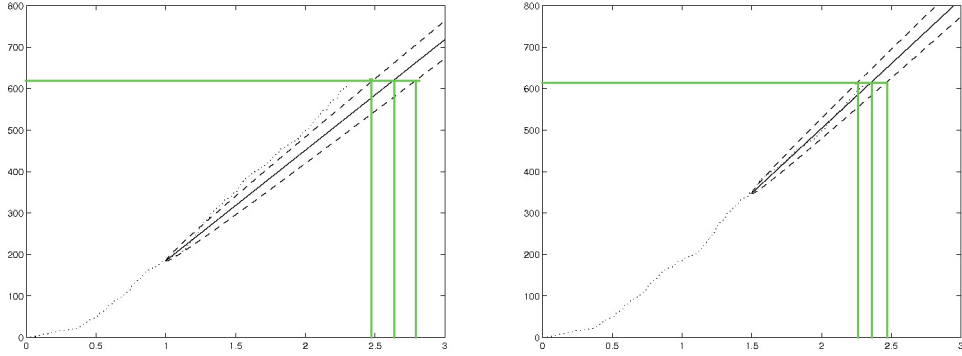


FIGURE 3.7 – Carte de contrôle du suivi de patients. Estimation du nombre moyen de patients recrutés à partir d'une étude intermédiaire à 1 an à gauche et 1.5 an à droite. En pointillés les intervalles de confiance à 95%. En vert, l'estimation de la durée de l'essai sachant que l'on souhaite recruter 610 patients.

Recrutement par centre/Ouverture et fermeture de centres

Les paramètres régissant la distribution des intensités de chaque centre sont nécessaires et calculés lors de l'application du modèle comme montré dans la Proposition 5, il est donc possible d'étudier la dynamique individuelle de chaque centre.

Il peut être intéressant d'étudier quels centres vont recruter au moins un patient pendant une durée T . En se plaçant dans le centre i et sachant la filtration $\mathcal{F}_{i,t_1}$ du processus $\mathcal{N}_i$ à l'instant t_1 , la probabilité pour que ce processus recrute un nombre $\tilde{n}$ de patients entre t_1 et $t > t_1$ est donnée par :

$$\mathbb{P}[\mathcal{N}_i(t) - k_i > \tilde{n} | \mathcal{F}_{i,t_1}] = 1 - \sum_{k=0}^{\tilde{n}-k_i-1} \frac{1}{k!} \int_{\mathbb{R}} \left(\int_{t_1}^t x dx \right)^k e^{-\left(\int_{t_1}^t x dx \right)} p_{\lambda_i}^{t_1}(x) dx,$$

Cette valeur peut se calculer numériquement et permet ainsi d'ajouter à l'étude une donnée de contrôle sur les centres investigateurs de l'étude. Ce calcul nous permet en effet de hiérarchiser les centres en fonction de leur recrutement futur prédit, et donc de pouvoir considérer un ajustement du nombre de centres en fonction d'objectif de temps de recrutement, de coût.

Ainsi dans [46] les auteurs utilisent des données réelles fournies par l'Unité Mixte de Recherche 1027 de l'Inserm. L'essai clinique comprenait 77 centres avait pour but de recruter 610 patients en 3 ans. En réalité, l'inclusion de patients s'est arrêté au bout

de 2,31 années. La particularité de ce jeu de données est que la date d'ouverture des centres n'a pas été fournie, en revanche, les dates d'inclusion de chaque patient étaient connues. Ainsi le modèle Poisson-gamma uniforme a été utilisé.

Le fait que la durée réelle est plus petite que la durée cible permet de déterminer (en prenant les centres recrutant le plus) le nombre de centres minimum nécessaire pour recruter les 610 patients en 3 ans.

En effet la flexibilité du modèle présenté permet de recalculer facilement la loi du temps d'inclusion si l'on décide de fermer ou ouvrir un centre. Si l'on ferme le centre j , l'intensité globale vaut $A_2 = \sum_{\substack{i=1 \\ i \neq j}}^C \lambda_i$, tandis que si l'on ouvre un nouveau centre, elle vaut $A_3 = \lambda + \sum_{i=1}^C \lambda_i$ où λ est l'intensité du nouveau centre.

Comme aucune information n'est disponible pour le nouveau centre, on suppose que son intensité λ suit la loi non conditionnée du modèle Poisson-gamma uniforme.

Ainsi, pour approcher la nouvelle intensité globale A par une loi Gamma, il suffit d'appliquer le théorème 3 avec :

$$\begin{aligned} \text{— si l'on ferme le centre } j, m &= \sum_{\substack{i=1 \\ i \neq j}}^C \mathbb{E}[\lambda_i | k_i], \sigma^2 = \sum_{\substack{i=1 \\ i \neq j}}^C \mathbb{V}[\lambda_i | k_i], \\ \text{— si l'on ouvre un centre, } m &= \mathbb{E}[\lambda] + \sum_{i=1}^C \mathbb{E}[\lambda_i | k_i], \sigma^2 = \mathbb{V}[\lambda] + \sum_{i=1}^C \mathbb{V}[\lambda_i | k_i]. \end{aligned}$$

En appliquant cette méthode aux données fournies par l'Inserm à $t_1 = 1.5$, l'estimateur du temps moyen de recrutement est $\hat{T} \simeq 2.34$ années et la probabilité de finir l'essai dans les 3 ans est très proche de 1. En ne gardant ouverts que les 39 centres (sur 77) ayant recruté au moins 3 patients, l'estimateur du temps moyen de recrutement est 2.73 années et la probabilité de finir l'essai en moins de 3 ans est de 99.5%.

Dans le cas où la prédiction $\tilde{T}$ faite en t_1 est trop grande par rapport à une durée maximum tolérée T_r pour la phase de recrutement de l'essai clinique, il peut être intéressant de savoir combien de nouveaux centres il est nécessaire d'inclure pour que la prédiction soit inférieure à T_r .

Théorème 6 *Pour plus de simplicité, supposons que pour tout $i \in \{1, \dots, C\}$, $\tau_i = \tau$ et que les nouveaux centres sont ouverts avec un délai de durée d . En supposant également les paramètres (α, μ) connus, alors pour finir le recrutement avant T_r avec une probabilité $1 - \delta$, il faut ouvrir un minimum de M centres où M est le plus petit entier vérifiant l'inégalité :*

$$M \geq \frac{A + Bz_{1-\delta}^2/2 + z_{1-\delta}\sqrt{AB + Q + B^2z_{1-\delta}^2/4}}{\alpha B},$$

où $A = \tilde{n} - (\alpha C + N_{t_1})(T_r - t_1)/(\alpha/\mu + \tau)$, $B = \mu(T_r - t_1 - d)/\alpha$, $Q = \tilde{n} + (\alpha C + N_{t_1})(T_r - t_1)^2/(\alpha/\mu + \tau)^2$ et $z_{1-\delta}$ est le quantile d'ordre $(1 - \delta)$ d'une distribution normale centrée réduite.

Preuve Notons d'abord que après l'instant t_1 , les C centres continuent de recruter avec une intensité globale Λ , où l'on utilise à la place l'estimation $\hat{\Lambda} = \sum_{i=1}^C \text{Gamma}(\alpha + k_i, \beta + \tau) = \text{Gamma}(C\alpha + K, \beta + \tau)$, où K est le nombre total de patients recrutés jusqu'à t_1 . En ajoutant M centres avec le même délai d , au temps $t_1 + d$ ces nouveaux centres ajoutent à l'intensité globale le terme $\Lambda_+ = \text{Gamma}(\alpha M, \beta)$. Ainsi le nombre de patients recrutés dans l'intervalle $[t_1, T]$ peut être représenté par $X = \Pi_{\hat{\Lambda}}(T - t_1) = \Pi_{\Lambda_+}(T - t_1 - d)$, où $\Pi_{\lambda}(t)$ représente un processus de Poisson d'intensité λ pris au temps t .

Le recrutement sera fini à temps dès lors que $X \geq n - K$. Il faut donc trouver M tel que $\mathbb{P}[X \geq n - K] \geq 1 - \delta$. En utilisant l'approximation des distribution de Poisson et gamma ainsi que $\mathbb{P}[T(n, C) \leq t] = \mathbb{P}(\text{Gamma}(n, 1) \leq \sum_{i=1}^C (t - u)\lambda_i \chi(t \geq u))$, nous obtenons le résultat souhaité.

Une étude par simulation [10] montre un cas où le nombre de centres est insuffisant pour atteindre les n patients en un temps donné T . Les paramètres de la simulation sont les suivants : $n = 400$, $C = 70$, $T = 365$ jours, tout les centres sont ouverts après au plus 5 mois. La perturbation vient du fait que la moitié des sites sont fermés deux mois avant la fin du recrutement. La Figure 3.8 montre les résultats de cette analyse.

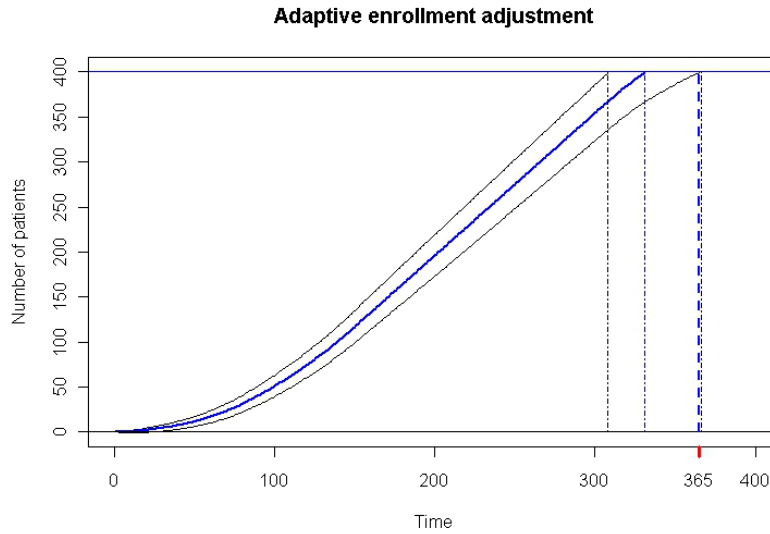
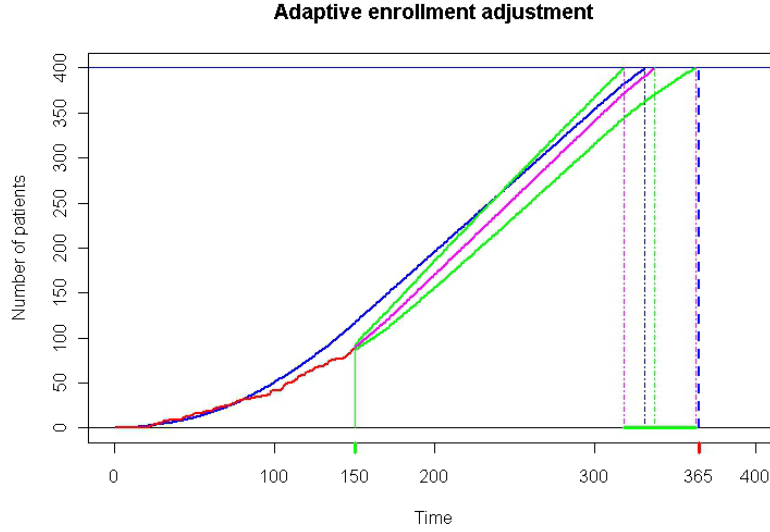


FIGURE 3.8 – Prévission pré-recrutement de la dynamique de recrutement avec intervalle de confiance à 90%.

Un premier temps d'analyse $t_1 = 150$ jours est considéré, 88 patients sont alors déjà recrutés à cet instant, le modèle prédit à cet instant un recrutement plus lent que ce qui est attendu, cependant cet aspect peut être relié à la phase d'apprentissage présentée au paragraphe 2.1, cependant ce phénomène sera substituer par la fermeture des centres, un ajustement du recrutement et donc nécessaire si le recrutement doit finir au temps T .

Ainsi en prenant en compte les centres qui vont fermés, le modèle prédit un recrutement finissant à temps si au moins 22 nouveaux centres sont ajoutés à partir de $t_1 = 150$ jours.


 FIGURE 3.9 – Courbe magenta : processus ajusté à $t_1 = 150$ avec ajout de 22 centres.

Patient dans l'essai avec considération du temps de suivi

Considérons un essai clinique en cours de réalisation comportant C centres. Une fois la randomisation d'un patient effectuée, celui-ci se retrouve affecté à un traitement k . À partir de ce moment, il va être suivi lors de plusieurs visites pendant un certain temps que nous noterons L et considérons constant pour tous les patients. Les phases de recrutement et de suivi se superposent, en effet certains patients recrutés en début d'étude peuvent être en période de suivi tandis que d'autres patients sont encore à recruter. Nous considérons également le phénomène de sortie d'étude comme au paragraphe 2.3. Ici le k -ème patient recruté par le centre i peut sortir de l'étude selon une intensité $\mu_{k,i}$.

Nous allons définir un processus général qui peut être défini à partir d'un processus de recrutement et qui va nous servir à modéliser le suivi des patients. Pour le centre i , définissons les temps d'arrivée $t_{1,i} \leq t_{2,i} \leq \dots$. On associe à chaque $t_{k,i}$ un processus aléatoire :

$$\xi_{k,i}(t) = 1, 0 \leq t \leq T_{k,i}; \xi_{k,i}(t) = 0, t \geq T_{k,i},$$

où $T_{k,i} = \min(\mathbb{E}[\mu_{k,i}], L)$.

Le processus $\xi_{k,i}$ défini ainsi pour le k -ème patient recruté par le centre i s'il est encore présent dans l'étude ou non.

Ainsi en considérant le processus $Z_i(t) = \sum_k \xi_{k,i}(t - t_{k,i})$, nous obtenons le processus du nombre de patients encore présent dans l'étude pour le centre i à l'instant t , et $Z(t) = \sum_{i=1}^C Z_i(t)$ le nombre total de patients encore présent dans l'étude à l'instant t .

Il est alors possible de calculer les moyennes et variances de ce processus pour prédire le nombre de patients présents dans l'étude pour tout instant t .

La figure 3.10 montre une étude par simulation réalisée où 80 centres sont considérés, $n = 400$ et $L = 4$ mois, $\mu = 0.001$ est constant pour tous les centres pour une durée de recrutement totale de 1 an.

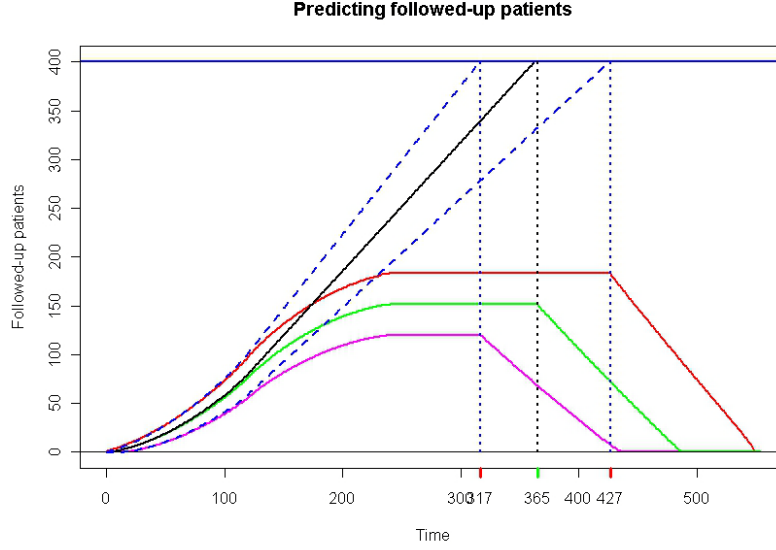


FIGURE 3.10 – Nombre de patients suivis au cours du recrutement. Courbe noir (pointillés bleus) : moyenne du processus prédictif de recrutement (intervalles de confiance à 90%). Courbe verte : moyenne du processus $Z(t)$, Courbe rouge (magenta) : limites prédictives supérieure de $Z(t)$ à 90 % (resp. inférieure)

Ainsi dans cette étude jamais plus de 150 patients sont en même temps en période de suivi. Cette donnée peut servir pour de nombreuses problématiques comme la taille des stocks de traitement nécessaires, le nombre de personnel de suivi.

3.3 Au suivi dynamique du recrutement d'un essai clinique avec outcome longitudinal

Lorsque le critère de jugement est longitudinal, des données de survie par exemple, alors à la dynamique de recrutement s'ajoute la dynamique de survie, découlant sur des modèles hiérarchiques. Ces deux processus étant indépendants, l'utilisation hiérarchique des modèles Bayésiens ne pose pas de problème majeur. Ce modèle permet alors de calculer ou au moins de simuler le temps d'atteinte non pas d'un nombre fixé de patients mais d'un nombre fixé d'événements que l'on souhaite observer (event-driven design). De nombreuses publications utilisent des processus de Poisson pour la modélisation d'apparition d'évènement, nous pouvons citer entre autre les travaux de Bagiella et al. [12], Donovan et al. [27] dans le cas d'étude de données de survie exponentielles qui considère des estimateurs ponctuels pour la prédiction du temps de l'essai. Citons également les travaux de Ying et al. [64] qui utilisent des hypothèses d'estimations non paramétriques ainsi que Ying et al. [63] où est considéré le temps de survie de Weibull. Cependant pour évaluer également un intervalle de confiance, nous utiliserons les résultats d'une étude par simulation Monte-Carlo [5]. Il est ainsi possible d'estimer non pas la durée de la phase de recrutement mais la durée total de l'étude.

3.3.1 Outcome : données de survie

Considérons premièrement le cas de l'observation d'un seul évènement que nous notons E . Soient τ_E le temps nécessaire pour que l'évènement E se produise (τ_E sera considéré par la suite comme une variable exponentielle d'intensité μ), et $p_E(x) = \mathbb{P}[\tau_E \leq x]$, $x \geq 0$ la probabilité que l'évènement E survienne. Considérons maintenant un processus de recrutement pour un essai clinique modélisé par un processus de Poisson d'intensité $\lambda(u)$ pouvant être aléatoire. Soient Π_a une variable de Poisson d'intensité a et $\Pi_a(t)$ son processus.

Lemme 2 *La prédiction du nombre d'évènement apparaissant dans l'intervalle $[0, t]$ pour les nouveaux patients venant d'être recrutés possède une distribution de Poisson d'intensité $\int_0^t \lambda(u)p_E(t-u) du$*

Preuve Sur un intervalle court $[u, u+h]$, le nombre de patients recrutés est approché par une variable de Poisson $\Pi_{\lambda(u)h}$. Pour chaque nouveaux patients recrutés dans cet intervalle, l'évènement E peut apparaître dans l'intervalle de temps restant $[u, t]$ avec probabilité $p_E(t_u)$. Ainsi le nombre d'évènements dans l'intervalle $[u, t]$ pour les patients recrutés dans l'intervalle $[u, u+h]$ est approché par une distribution de Poisson d'intensité $\lambda(u)p_E(t_u)h$. Sachant que la somme de variables de Poisson indépendantes est une variable de Poisson où la somme des intensités donne son intensité, en faisant tendre h vers 0 nous prouvons le Lemme.

Dans le cas d'un recrutement multicentrique faisant intervenir C centres possédant chacun une intensité de recrutement $\lambda_i(u)$, la prédiction du nombre total d'évènement $Z(t, E)$ apparaissant dans l'intervalle $[0, t]$ possède une distribution de Poisson doublement stochastique d'intensité cumulée $\Sigma(t, E)$, où :

$$\Sigma(t, E) = \sum_{i=1}^C \int_0^t \lambda_i(u)p_E(t_u) du \quad (3.4)$$

Nous nous plaçons maintenant dans le cas où les processus de Poisson possède des intensités aléatoires selon une distribution Gamma.

Prédiction du nombre d'évènement

On suppose qu'au temps t_1 C centres sont actifs et que le recrutement sera fini au temps $t_0 + T$. Les données de recrutement $\{k_i, \tau_i\}_{i=1, \dots, C}$ (k_i : nombre de patients recrutés par le centre i , τ_i : temps d'activité du centre i jusqu'à t_1) sont supposées données.

Pour simplifier les notations posons $t_1 = 0$ et :

$$q(t, E, T) = \int_0^{\min(t, T)} p_E(t-u) du, \quad t \geq 0. \quad (3.5)$$

Théorème 7 [5] *La variable $Z(t, E)$ possède une distribution doublement stochastique d'intensité $q(t, T, E)\Lambda$ et $\mathbb{E}[Z(t, E)] = q(t, T, E)A$, $\mathbb{V}[Z(t, E)] = q(t, T, E)\mathbf{A} + q^2(t, T, E)\mathbf{B}$, où $\mathbf{A}$ et $\mathbf{B}$ sont données par le théorème 2.*

Preuve En utilisant les formules de la preuve du théorème 5, le résultat découle directement en considérant le Lemme 2 avec comme intensité cumulée $\Lambda = \sum_{i=1}^C \lambda_i$.

Plaçons nous dans le contexte de l'étude menée par Anisimov dans [5], à savoir 227 patients déjà recrutés, 131 centres actifs, 32 centres qui vont être ajoutés. AU temps d'analyse $n_0 = 175$ patients sont à risques et $n_E = 52$ évènements se sont produits, et $\mu = 0.001482$.

La prédiction du nombre total d'évènement $Y(t)$ peut être représenté comme :

$$Y(t) = n_E + \text{Bin}(n_0, p_\mu(t)) + Z(t, E),$$

où $p_\mu(t) = 1 - \exp^{-\mu t}$.

On peut alors calculer la moyenne et la variance de $Y(t)$ servant pour la prédiction et l'intervalle de confiance.

Les résultats dans le cas de 2 scénarios (Scénario 1 : arrêt du recrutement après 1 ans, Scénario 2 : arrêt du recrutement après 2 ans) sont présenté dans la figure 3.11.

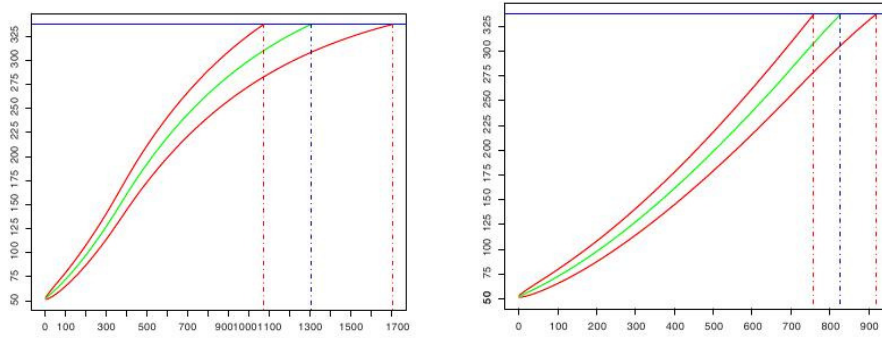


FIGURE 3.11 – Nombre d'évènements observés en fonction du temps ainsi que la durée moyenne de l'essai et son intervalle de prédiction à 90% pour observer 337 évènements. A gauche le recrutement s'arrête au bout d'un an, la durée moyenne est de 1304 (1072, 1705) jours. A droite le recrutement s'arrête au bout de 2 ans, la durée moyenne est de 757 (698, 828) jours.

3.3.2 Essai clinique avec plusieurs visites : multi-états

Dans la même veine on peut envisager une dynamique de recrutement Poisson-gamma et une fois le patient recruté il entre dans une chaîne de Markov à espace d'états fini. Le modèle permet alors, à partir des données recueillies à un instant intermédiaire de l'étude, d'estimer le nombre de sujets se trouvant dans un certain état. En supposant que les probabilités d'être dans un certain état à un instant donné soient connues, il est possible à partir des estimations du modèle de recrutement de créer un modèle, se basant sur des chaînes de Markov finies, qui donne le nombre de sujets dans un état en fonction du temps (modèle qui est en effet dépendant du recrutement dans le sens où tous les patients ne passent pas les visites au même moment, certains ont fini leurs visites avant que d'autres sujets n'est été recrutés ou screenés). Ces données peuvent par exemple servir à redéfinir le nombre de sujets nécessaire dans le cas où l'état d'intérêt a été grandement sous ou sur-estimé, et donc de redéfinir également comme expliqué précédemment le nombre de centres nécessaires. Ces techniques devaient être appliquées sur des données disponibles à l'unité 1027 malheureusement le trop faible nombre de centres n'a pas permis une étude conclusive.

Une présentation du modèle est tout de même donnée par la suite.

Modèle de markov fini pour l'étude de perte de patients lors des visites

Dans cette partie nous nous intéresserons à la survenue de 3 évènements, R en rechute et non décédé, D décédé et L perdu de vue lors de la phase de suivi. On divise l'ensemble des patients en six groupes distincts :

- groupe O : k_O patients recrutés et à risques.
- groupe R : k_R patients recrutés avec rechute encore suivi.
- groupe D_1 : k_{D_1} patients décédés avec la mort comme premier évènement survenu.
- groupe D_2 : k_{D_2} patients décédés où le décès est survenu après la rechute.
- groupe L_1 : k_{L_1} patients perdus de vue lors de la phase de suivi avant la survenue d'un évènement.
- groupe L_2 : k_{L_2} patients perdus de vue lors de la phase de suivi après la survenue de la rechute.

On suppose de nouveau que la survenue des évènements possède une distribution exponentielle. On considère six intensités de survenues des évènements, à savoir μ_R l'intensité de la rechute, μ_{D_1} l'intensité du décès dans le groupe O , μ_{D_2} l'intensité de décès dans le groupe R , μ_{L_1} celle de l'évènement L dans le groupe O et μ_{L_2} l'intensité de l'évènement L dans le groupe R .

Le but est maintenant de prédire dans le temps le nombre d'évènements pour chaque groupe.

Le comportement de chaque patients peut être modélisé par une chaîne de Markov à temps continu à six états $\Omega = \{O, R, D_1, D_2, L_1, L_2\}$. Les probabilités de transition peuvent alors être calculées comme montré dans [5, Annexe A.2].

3.4 A la modélisation du coût¹

3.4.1 Introduction

L'évaluation du coût d'un essai clinique est un problème complexe. Une étude durant plusieurs années et s'appuyant sur plusieurs dizaines de centres implique une logistique très importante, qui engendre des dépenses tout aussi importantes. Néanmoins, il est difficile d'évaluer précisément ce coût, car une multitude de paramètres entrent en jeu : coût de fabrication et d'acheminement des médicaments, coût du personnel, de prise en charge des patients, du management de l'étude...

Nous supposons que l'on peut néanmoins catégoriser ces coûts de la manière suivante :

- **coût fixe par patient screené**,
- **coût fixe par patient randomisé** incluant entre autres le coût du traitement,
- **coût par patient randomisé dépendant du temps passé dans l'essai** : si les patients ne restent pas tous le même temps dans l'essai (à cause par exemple

1. Ce travail est le résultat d'une collaboration avec Guillaume Mijoule, Vladimir Anisimov et Nicolas Savy, et est présenté dans les actes du 7th International Workshop on Statistical Simulation [45].

de l'évolution de leur état de santé ; ce sera en particulier le cas dans le cadre des données de survie),

- **coût fixe d'un centre** qui peut inclure le coût de mise en route du centre, de la fourniture en médicaments (en général, la quantité de médicaments fournie à un centre est déterminée au début de l'essai et ne dépend donc pas du processus d'inclusion du centre),
- **coût d'un centre proportionnel à sa durée d'activité** incluant une partie du coût du personnel (par exemple, un médecin est requis en permanence dans un centre ouvert, que le centre recrute ou non), les coûts en énergie, etc.

Le coût de l'essai pour le centre i à l'instant $t \geq 0$ peut donc s'écrire :

$$C_i(t) = C_1 K_i(t) + C_2 N_i(t) + \sum_{0 \leq T_n^i \leq t} g(t, T_n^i) + F_i + G_i t, \quad (3.6)$$

où $(T_n^i)_{n \geq 0}$ sont les instants de randomisation dans le centre i , $g(t, s)$ représente le coût à l'instant t d'un patient randomisé à l'instant s , F_i est le coût fixe du centre, G_i le coût par année d'activité du centre, C_1 le coût fixe d'un patient randomisé, C_2 le coût de screening d'un patient. On fait les hypothèses suivantes sur g :

- (H1) $g : \mathbb{R}_+ \times \mathbb{R}_+ \rightarrow \mathbb{R}$ est mesurable,
- (H2) g triangulaire au sens où $g(t, s) = 0$ si $t < s$,
- (H3) pour tout $t \geq 0$, $g(t, \cdot)$ est continue sur $[0, t]$.

Notons que la somme dans (3.6) peut s'écrire comme une intégrale contre le processus K_i :

$$\sum_{0 \leq T_n^i \leq t} g(t, T_n^i) = \int_0^t g(t, s) dK_i(s).$$

Si K_i est un processus de Poisson d'intensité homogène, le processus $t \mapsto \left(\int_0^t g(t, s) dK_i(s) \right)_{t \geq 0}$ est connu sous le nom de processus de Poisson filtré [21, 22, 54, 55].

Le coût global de l'essai à l'instant t est alors défini par

$$C(t) = \sum_{i=1}^C C_i(t).$$

L'inclusion de patients s'arrête dès que le processus $K = \sum_{i=1}^C K_i$ atteint K_f . En posant

$$\tau = \left\{ \inf_{t \geq 0} : K(t) = K_f \right\}$$

le coût total de l'essai est alors $C(\tau)$. On en déduit alors le coût moyen de l'essai, noté $\mathcal{C}$:

coût moyen de l'essai On appelle coût moyen de l'essai la quantité

$$\mathcal{C} = \mathbb{E}[C(\tau)] = C_1 \mathbb{E}[K(\tau)] + C_2 \mathbb{E}[N(\tau)] + \mathbb{E} \left[\int_0^\tau g(\tau, s) dK(s) \right] + \sum_{i=1}^C F_i + \sum_{i=1}^C G_i \mathbb{E}[\tau]. \quad (3.7)$$

L'intégrale précédente s'entend au sens de Stieljès : $\int_0^\tau g(\tau, s) dK(s) = \sum_{0 \leq T_n^i \leq \tau} g(\tau, T_n^i)$, $(T_n^i)_{n \geq 0}$ étant les instants de randomisation dans le centre i . Notons que le processus $t \mapsto \int_0^t g(t, s) dK(s)$ n'étant pas une semi-martingale [22], nous ne pouvons pas calculer l'espérance par des techniques de martingale.

On se place dans le cadre du modèle 2 du paragraphe 2.3. On rappelle que le processus N_i est un processus de Poisson d'intensité λ_i , et K_i est un processus de Poisson de paramètre $p_i \lambda_i$, où λ_i suit une loi Gamma de paramètres (α, β) et p_i suit une loi Beta de paramètres (ψ_1, ψ_2) .

3.4.2 Préliminaires

Reprenons la formule (3.7). Supposons d'abord $(\lambda_i, p_i)_{1 \leq i \leq C}$ connus. Sachant $(p_i, \lambda_i)_{1 \leq i \leq C}$, K est un processus de Poisson d'intensité $\Lambda = \sum_{i=1}^C p_i \lambda_i$. Le processus N est alors représenté par $N \equiv K + K_2$ où K_2 est un processus de Poisson indépendant de K , d'intensité $\Lambda_2 = \sum_{i=1}^C (1 - p_i) \lambda_i$. On se ramène alors au problème suivant :

Soient K et K_2 deux processus de Poisson indépendants d'intensités respectives Λ et Λ_2 . Posons

$$C(t) = C'_1 K(t) + C_2 K_2(t) + \int_0^t g(t, s) dK(s) + Gt + F, \quad t \geq 0$$

où g vérifie les hypothèses (H1), (H2) et (H3). Dans le cadre des essais cliniques, on aura $F = \sum_{i=1}^C F_i$, $G = \sum_{i=1}^C G_i$, et $C'_1 = C_1 + C_2$.

Soit $\tau = \inf\{t \geq 0 : K(t) \geq K_f\}$. Le coût total de l'essai est donc

$$C(\tau) = C'_1 K_f + C_2 K_2(\tau) + \int_0^\tau g(\tau, s) dK(s) + G\tau + F. \quad (3.8)$$

Si K et K_2 sont deux processus de Poisson indépendants d'intensités respectives Λ et Λ_2 , et si $\tau = \inf\{t \geq 0 : K(t) \geq K_f\}$, alors le coût moyen défini dans (3.7) vaut

$$\begin{aligned} \mathbb{E}[C(\tau)] &= C'_1 K_f + C_2 K_f \frac{\Lambda_2}{\Lambda} + \int_0^{+\infty} g(t, t) p_\tau(dt) + (K_f - 1) \int_0^{+\infty} \int_0^t g(t, s) ds \frac{p_\tau(dt)}{t} \\ &\quad + G \frac{K_f}{\Lambda} + F. \end{aligned}$$

La démonstration repose sur les lemmes suivants.

Lemme 3 *Si K est un processus de Poisson d'intensité Λ et $\tau = \inf\{t \geq 0 : K(t) \geq K_f\}$, alors*

$$\mathbb{E}[\tau] = \frac{K_f}{\Lambda}.$$

Preuve D'après le théorème 2, τ suit une loi Gamma de paramètres (K_f, Λ) , d'où le résultat.

Lemme 4 *Soient K et K_2 deux processus de Poisson indépendants, d'intensités respectives Λ et Λ_2 . Soit $\tau = \inf\{t \geq 0 : K(t) \geq K_f\}$. Alors*

$$\mathbb{E}[K_2(\tau)] = K_f \frac{\Lambda_2}{\Lambda}. \quad (3.9)$$

Preuve Comme la variable aléatoire τ est indépendant de K_2 ,

$$\mathbb{E}[K_2(\tau)] = \mathbb{E}[\mathbb{E}[K_2(\tau) | \tau]] = \mathbb{E}[\Lambda_2 \tau] = \Lambda_2 \mathbb{E}[\tau] = K_f \frac{\Lambda_2}{\Lambda}.$$

Lemme 5 Posons $n = K_f - 1$. Sachant τ , les instants de saut $T_1 \leq T_2 \leq \dots \leq T_n$ de K sur $[0, \tau]$ ont même loi qu'un n -uplet de loi uniforme sur $[0, \tau]^n$ réordonné par ordre croissant.

Preuve En effet, le $(n+1)$ -uplet $(T_1, \dots, T_n, T_{n+1} = \tau)$ admet une densité sur $\mathbb{R}_+^{n+1}$ qui est

$$\begin{aligned} p_{T_1, \dots, T_{n+1}}(t_1, \dots, t_{n+1}) &= \chi(0 \leq t_1 \leq \dots \leq t_{n+1}) \prod_{i=1}^{n+1} \Lambda e^{-\Lambda(t_i - t_{i-1})} \\ &= \chi(0 \leq t_1 \leq \dots \leq t_{n+1}) \Lambda^{n+1} e^{-\Lambda t_{n+1}}, \end{aligned}$$

où par convention $t_0 = 0$. $\tau = T_{n+1}$ suit une loi $\Gamma(K_f, \Lambda)$. Notons p_τ sa densité, et rappelons que

$$p_\tau(dt) = \chi(t > 0) \frac{\Lambda^{K_f}}{(K_f - 1)!} e^{-\Lambda t} t^{K_f - 1} dt. \quad (3.10)$$

Sachant $T_{n+1} = t$, la densité de $(T_1, \dots, T_n)$ s'écrit alors

$$\begin{aligned} p_{T_1, \dots, T_n | T_{n+1} = t}(t_1, \dots, t_n) &= \frac{\Lambda^{n+1} e^{-\Lambda t}}{p_\tau(t)} \chi(0 \leq t_1 \leq \dots \leq t_n \leq t) \\ &= \frac{n!}{t^n} \chi(0 \leq t_1 \leq \dots \leq t_n \leq t). \end{aligned}$$

On en déduit le corollaire :

Corollaire 2 Sachant τ , pour tout $s < \tau$, $K(s)$ suit une loi binomiale $\mathcal{B}(K_f - 1, \frac{s}{\tau})$.

Preuve C'est une conséquence directe de la proposition précédente.

Lemme 6 Supposons que g satisfait les hypothèses (H1) à (H3). Alors

$$\mathbb{E} \left[\int_0^\tau g(\tau, s) dK(s) \right] = \int_0^{+\infty} g(t, t) p_\tau(dt) + (K_f - 1) \int_0^{+\infty} \int_0^t g(t, s) ds \frac{p_\tau(dt)}{t} \quad (3.11)$$

$$= \mathbb{E}[g(\tau, \tau)] + (K_f - 1) \mathbb{E}[G(\tau)], \quad (3.12)$$

où pour tout $t > 0$, $G(t) = \frac{1}{t} \int_0^t g(t, s) ds$.

Preuve On a

$$\mathbb{E} \left[\int_0^\tau g(\tau, s) dK(s) \right] = \int_0^{+\infty} \mathbb{E} \left[\int_0^t g(t, s) dK(s) \mid \tau = t \right] p_\tau(dt),$$

Supposons d'abord que pour tout $t > 0$, la restriction de $s \mapsto g(t, s)$ à $[0, t]$ est dérivable. Notons $\partial_2 g(t, \cdot)$ cette dérivée. Supposons en outre que $\forall t > 0, \sup_{0 \leq s \leq t} |\partial_2 g(t, s)| < +\infty$.

Par une intégration par parties, on a

$$\int_0^t g(t, s) dK(s) = [g(t, s) K(s)]_0^t - \int_0^t \partial_2 g(t, s) K(s) ds,$$

soit

$$\int_0^t g(t, s) dK(s) = g(t, t)K(t) - \int_0^t \partial_2 g(t, s) K(s) ds. \quad (3.13)$$

On en déduit que

$$\mathbb{E} \left[\int_0^t g(t, s) dK(s) \mid \tau = t \right] = g(t, t)K_f - \mathbb{E} \left[\int_0^t \partial_2 g(t, s) K(s) ds \mid \tau = t \right].$$

Sachant $\{\tau = t\}$, nous pouvons majorer $|\partial_2 g(t, s) K(s)| \leq K_f \sup_{0 \leq s \leq t} |\partial_2 g(t, s)|$, le théorème de Fubini s'applique et grâce au corollaire 2,

$$\mathbb{E} \left[\int_0^t \partial_2 g(t, s) K(s) ds \mid \tau = t \right] = (K_f - 1) \int_0^t \partial_2 g(t, s) \frac{s}{t} ds,$$

d'où l'on déduit (3.12) par une intégration par parties.

Supposons à présent que pour tout $t > 0$, $g(t, \cdot)$ est continue sur $[0, t]$. Alors pour tout $\epsilon > 0$, il existe une fonction $\phi(t, \cdot) \mathcal{C}^1$ sur $[0, t]$ telle que $\sup_{0 \leq s \leq t} |\phi(t, \cdot) - g(t, \cdot)| \leq \epsilon$, et

$$\begin{aligned} & \left| \mathbb{E} \left[\int_0^\tau g(\tau, s) dK(s) \right] - \int_0^{+\infty} g(t, t) p_\tau(dt) - (K_f - 1) \int_0^{+\infty} \int_0^t g(t, s) ds \frac{p_\tau(dt)}{t} \right| \\ &= \left| \mathbb{E} \left[\int_0^\tau (g - \phi)(\tau, s) dK(s) \right] - \int_0^{+\infty} (g - \phi)(t, t) p_\tau(dt) \right. \\ & \quad \left. - (K_f - 1) \int_0^{+\infty} \int_0^t (g - \phi)(t, s) ds \frac{p_\tau(dt)}{t} \right| \\ &\leq K_f \epsilon + (K_f - 1) \int_0^{+\infty} t \frac{p_\tau(dt)}{t} \\ &\leq (2K_f + 1). \end{aligned}$$

étant choisi arbitrairement, on en déduit le résultat.

3.4.3 Application aux essais cliniques multicentriques

On suppose que la fonction de coût g s'écrit

$$g(t, s) = C_3(t - s)\chi(t \geq s), \quad (3.14)$$

où $C_3 > 0$. Ceci revient à dire qu'un patient randomisé représente, en plus d'un coût fixe C_1 , un coût proportionnel au temps qu'il passe dans l'essai.

Paramètres $(\lambda_i)_{1 \leq i \leq C}$ et $(p_i)_{1 \leq i \leq C}$ connus

Dans un premier temps, supposons les intensités d'inclusion $(\lambda_i)_{1 \leq i \leq C}$ et les probabilités de randomisation $(p_i)_{1 \leq i \leq C}$ connues.

L'application du théorème 3.4.2 nous donne le coût moyen :

$$C(\tau) = C'_1 K_f + C_2 K_f \frac{\Lambda_2}{\Lambda} + \frac{1}{2} C_3 (K_f - 1) \frac{K_f}{\Lambda} + G \frac{K_f}{\Lambda} + F, \quad (3.15)$$

où

$$\begin{aligned}\Lambda &= \sum_{i=1}^C p_i \lambda_i, \\ \Lambda_2 &= \sum_{i=1}^C (1 - p_i) \lambda_i, \\ F &= \sum_{i=1}^C F_i, \\ G &= \sum_{i=1}^C G_i.\end{aligned}$$

A-t-on intérêt à fermer le centre j ? La proposition suivante permet de répondre à cette question. La fermeture du centre j entraîne une variation du coût moyen de l'essai égale à

$$\Delta \mathcal{C}_j = K_f \frac{\lambda_j}{\Lambda(\Lambda - p_j \lambda_j)} \left[C_2(p_j \Lambda_2 - (1 - p_j)\Lambda) + \frac{1}{2} C_3(K_f - 1)p_j + p_j G - G_j \frac{\Lambda}{\lambda_j} \right] - F_j. \quad (3.16)$$

Le fait de fermer un centre augmente certes le temps de recrutement, donc le coût de chacun des centres non fermés, cette expression permet ainsi de répondre à la question : les coûts ajoutés des centres non fermés par l'augmentation du temps de recrutement du fait de la fermeture du centre j sont-ils au moins compensés par les gains engendrés par cette fermeture. Il est rentable de fermer le centre j si cette quantité est négative.

Preuve La fermeture du centre j induit une variation de Λ , Λ_2 , F et G égales à

$$\begin{aligned}\Delta \Lambda &= -p_j \lambda_j, \\ \Delta \Lambda_2 &= -(1 - p_j) \lambda_j, \\ \Delta F &= -F_j, \\ \Delta G &= -G_j.\end{aligned}$$

Le coût moyen étant donné par l'équation (3.15), on en déduit le résultat.

Paramètres $(\lambda_i)_{1 \leq i \leq C}$ et $(p_i)_{1 \leq i \leq C}$ aléatoires

On se place ici dans le cadre du modèle 2 du paragraphe 2.3 : les intensités d'inclusion des centres $\{\lambda_1, \dots, \lambda_C\}$ sont indépendantes de loi Gamma de paramètres (α, β) , et les probabilités de randomisation $\{p_1, \dots, p_C\}$ sont indépendantes de loi Beta de paramètres (ψ_1, ψ_2) .

On fait une étude à un instant intermédiaire t_1 . On considère donc le **coût de l'étude à partir de l'instant t_1** , et non pas le coût global de l'étude. Notons donc K_f le nombre total de patients restant à randomiser après t_1 . Sachant l'information à t_1 , en notant n_i le nombre de patients recrutés et k_i le nombre de patients randomisés dans le centre i , les lois de λ_i et p_i sont données par (cf paragraphe 2.3) :

$$\begin{aligned}\lambda_i &\stackrel{\mathcal{L}}{=} \Gamma(\alpha + n_i, \beta + \tau_i), \\ p_i &\stackrel{\mathcal{L}}{=} \mathcal{B}(\psi_1 + k_i, \psi_2 + n_i - k_i).\end{aligned}$$

Dans ce cadre, la proposition 3.4.3 implique, en moyennant sur les intensités d'inclusion $(\lambda_i)_{1 \leq i \leq C}$ et probabilités de randomisation $(p_i)_{1 \leq i \leq C}$: La fermeture du centre j entraîne une variation du coût moyen de l'essai égale à

$$\Delta_j = K_f \mathbb{E} \left[\frac{\lambda_j}{\Lambda(\Lambda - p_j \lambda_j)} \left[C_2(p_j \Lambda_2 - (1 - p_j)\Lambda) + \frac{1}{2} C_3(K_f - 1)p_j + p_j G - G_j \frac{\Lambda}{\lambda_j} \right] \right] - F_j. \quad (3.17)$$

Chapitre 4

Conclusion et perspectives

Les modèles bayésiens utilisés par Anisimov [10] et leurs extensions et variantes [46, 47] ne cessent d'être de plus en plus utilisés lors de la planification et l'élaboration des phases de recrutement de la phase III d'un essai clinique de par leur bonne adéquation aux données réelles, comme le montrent les tests effectués sur les courbes d'occupation, et la qualité de la prédiction qu'ils permettent. Les garanties théoriques permettent d'établir des estimations fiables lorsque le nombre de centres investigateurs est assez élevé, et les données recueillies durant l'étude intermédiaire peuvent servir à quantifier l'erreur commise sur les estimateurs du modèle et donc sur la prédiction.

De nombreuses études par simulations ont permis de mettre à jour la robustesse du modèle malgré la non adéquation des données étudiées aux hypothèses du modèle que se soit en terme d'interruptions des dynamiques de recrutement, de délai dans l'initialisation des centres ou des phases d'apprentissage et de tarissement.

Il a été également vu que les modèles pouvaient servir de base à des modèles plus complexes prenant en compte la perte de patients en phase de screening, eux-mêmes servant de base pour un modèle de coût. De plus, ils peuvent être simplement mis en œuvre numériquement et utilisés pour un usage en routine. L'essai peut également être suivi durant toute la phase de recrutement pour apporter des indicateurs pertinents sur des problématiques très variées (suivi du nombre de patients, du nombre de patients par centre/région, du nombre de traitements nécessaires par région, du nombre d'évènement survenant à un instant t , ...) notamment par l'utilisation de modèles hiérarchiques. Une introduction à l'étude du suivi des visites dont les données suivent des chaînes de Markov est également présentée.

Finalement un modèle permettant l'incorporation de données historiques pour la faisabilité d'un nouvel essai semble encourageant de par sa précision de prédiction sur données réelles [47].

4.1 Perspectives

Les perspectives de recherches s'articulent suivant deux axes, un axe théorique dont l'objectif est de mettre en avant les propriétés des modèles introduits dans cette thèse et un axe appliqué dont l'objectif est de développer ces outils et de les faire connaître au monde médical actuellement en grand changement. Pour ce dernier point nous avons disposé du soutien de l'IRESF qui a financé, dans le cadre de l'Appel d'Offre "Soutien à la recherche mathématique et statistique appliquée à la cancérologie" le projet "Statistical Modelling for Patients recruitment" porté par l'Institut de Mathématiques de Toulouse et l'Unité INSERM 1027 en partenariat avec Vladimir Anisimov, du soutien

du programme CAPTOR ("Cancer Pharmacology of Toulouse-Oncopole and Region") qui porte sur un projet de recherche consacré à l'innovation, l'évaluation, la diffusion des médicaments anti-cancéreux ainsi qu'à la formation relative à ces domaines, et également du soutien du projet HATICE ("Healthy ageing through Internet counselling in the elderly") qui développe une plateforme d'intervention en ligne pour aider les personnes âgées à réduire leur risque de maladie cardiovasculaire et de démence en favorisant la collaboration de multiples équipes statistiques au sein de l'UE.

4.1.1 Perspectives appliquées

Données de survie

Dans le cas des données de survie, un nouveau problème apparaît : ce qui est utile de connaître n'est pas nécessairement le nombre de patients mais le nombre d'événements d'intérêt à la fin de l'essai. Avec l'utilisation d'un modèle Gamma-Poisson pour l'inclusion et un modèle exponentiel pour les données de survie, Anisimov [5] obtient de bonnes prédictions du nombre d'événements.

Une perspective serait de développer un modèle plus fin que le modèle précédent, dans lequel une chaîne de Markov en temps continu est associée à chaque patient et dont les états représentent les différentes étapes de la progression de la maladie durant l'essai (à risque, perdu de vue, dans l'étude, décédé,...). En estimant les probabilités de transition, il sera possible de prédire plus précisément le nombre et le type d'événements à venir. Une étude d'un modèle de chaînes de Markov cachées a été réalisé, mais sans la possibilité de pouvoir le tester sur donnée réelle.

Cet axe d'étude est depuis quelque temps particulièrement étudié, avec ajout de covariables sur les états des patients lors des visites, nous pouvons ainsi citer [58, 59] pour des modèles simples, ainsi que [23, 29, 41] sur des modèles mixtes ou avec filtrage, et finalement les différentes contributions de de R. M. Altman [1, 2, 42].

Covariables déterminant les taux d'inclusion

Le point clé pour une bonne prédiction du temps de recrutement est une estimation la plus précise possible des intensités de recrutement de chaque centre. La recherche des covariables pouvant influencer la valeur de ces intensités permettrait d'améliorer la précision de la prédiction.

Il pourrait également être intéressant de modéliser les perturbations sur la dynamique de recrutement en terme de covariables pouvant être modélisée selon la même méthode bayésienne utilisée pour les intensités (le choix de la distribution sera à considérer).

4.1.2 Perspectives théoriques

Les perspectives théoriques concernent essentiellement les processus de Cox.

En premier lieu l'étude de sensibilité pourrait être améliorée en cherchant à estimer la distance de Wasserstein entre deux processus de Cox, vus comme variables aléatoires dans l'espace des configurations (cf [20]). Des travaux analogues de Barbour [14] dans le cas Poisson et Deucrosefond et al [20] dans le cas de processus ponctuels sont également des pistes à explorer.

Il serait également très utile de considérer une intensité dépendante du temps, permettant ainsi une généralisation du modèle à de nombreux cas se rapprochant de cas réels.

Bibliographie

- [1] R. M. Altman. Mixed hidden markov models : an extension of the hidden markov model to the longitudinal data setting. *Journal of the American Statistical Association*, 102(477) :201–210, 2007. [104](#)
- [2] R. M. Altman and A. J. Petkau. Application of hidden markov models to multiple sclerosis lesion count data. *Statistics in Medicine*, 24(15) :2335–2344, 2005. [104](#)
- [3] V. Anisimov. Predictive modelling of recruitment and drug supply in multicenter clinical trials. In *Proceedings of the Joint Statistical Meeting, ASA*, pages 1248–1259, Washington, USA, August 2009. [81](#)
- [4] V. Anisimov. Effects of unstratified and centre-stratified randomization in multicentre clinical trials. *Pharmaceutical Statistics*, 10(1) :50–59, 2011. [23](#), [84](#), [86](#), [87](#)
- [5] V. Anisimov. Predictive event modelling in multicentre clinical trials with waiting time to response. *Pharmaceutical Statistics*, 10(6) :517–522, 2011. [23](#), [92](#), [93](#), [94](#), [95](#), [104](#)
- [6] V. Anisimov. Statistical modeling of clinical trials (recruitment and randomization). *Comm. Statist. Theory Methods*, 40(19-20) :3684–3699, 2011. [15](#), [19](#), [20](#), [23](#), [44](#), [46](#), [49](#), [68](#), [75](#)
- [7] V. Anisimov, D. Downing, and V. Fedorov. Recruitment in multicentre trials : prediction and adjustment. In *mODa 8-Advances in Model-Oriented Design and Analysis*, pages 1–8. Springer, 2007. [14](#), [19](#), [20](#)
- [8] V. Anisimov, D. Downing, and V. Fedorov. Recruitment in multicentre trials : prediction and adjustment. 8th International Workshop in model-oriented design and analysis, pages 1–8, Almagro, Spain, June 2007. Physica-Verlag/Springer, Heidelberg. [44](#)
- [9] V. Anisimov and V. Fedorov. Modeling of enrolment and estimation of parameters in multicentre trials. *GlaxoSmithKline Pharmaceuticals*, 2005. [66](#)
- [10] V. Anisimov and V. Fedorov. Modelling, prediction and adaptive adjustment of recruitment in multicentre trials. *Statistics in Medicine*, 26(27) :4958–4975, 2007. [7](#), [10](#), [12](#), [14](#), [17](#), [18](#), [20](#), [21](#), [22](#), [54](#), [57](#), [68](#), [90](#), [103](#)
- [11] V. Anisimov, G. Mijoule, and N. Savy. Statistical modelling of recruitment in multicentre clinical trials with patients’ dropout. *Statistics in Medicine*, 2016. In preparation. [23](#), [39](#)
- [12] E. Bagiella and D. Heitjan. Predicting analysis times in randomized clinical trials. *Statistics in medicine*, 20(14) :2055–2063, 2001. [92](#)
- [13] A. Bakhshi, S. Senn, and A. Phillips. Some issues in predicting patient recruitment in multi- centre clinical trials. *Statistics in Medicine*, 32(30) :5458–5468, 2013. [66](#), [67](#)

- [14] AD Barbour, T. C Brown, and A. Xia. Point processes in time and stein's method. *Stochastics : An International Journal of Probability and Stochastic Processes*, 65(1-2) :127–151, 1998. [104](#)
- [15] K. Barnard, L. Dent, and A. Cook. A systematic review of models to predict recruitment to multicentre trials. *BMC Medical Research Methodology*, 63(10), 2010. [7](#)
- [16] J. M. Bernardo and A. F. M. Smith. *Bayesian Theory*. John Wiley & Sons, Hoboken, NJ, USA, 2004. [42](#)
- [17] R. E. Carter. Application of stochastic processes to participant recruitment in clinical trials. *Controlled Clinical Trials*, 25(5) :429–436, 2004. [7](#)
- [18] R. E. Carter, S. C. Sonne, and K. T. Brady. Practical considerations for estimating clinical trial accrual periods : application to a multi-center effectiveness study. *BMC Medical Research Methodology*, 5(11) :1–5, 2005. [7](#)
- [19] L. D. Case and T. M. Morgan. Duration of accrual and follow-up for two-stage clinical trials. *Lifetime Data Analysis*, 7(1) :21–37, 2001. [6](#)
- [20] L. Decreusefond, A. Joulin, and N. Savy. Upper bounds on rubinstein distances on configuration spaces and applications. *arXiv preprint arXiv :0812.3221*, 2008. [104](#)
- [21] L. Decreusefond and N. Savy. Filtered Brownian motions as weak limit of filtered Poisson processes. *Bernoulli*, 11(2) :283–292, 2005. [96](#)
- [22] L. Decreusefond and N. Savy. Anticipative calculus with respect to filtered Poisson processes. *Ann. Inst. H. Poincaré Probab. Statist.*, 42(3) :343–372, 2006. [96](#), [97](#)
- [23] J. C Detilleux. The analysis of disease biomarker data using a mixed hidden markov model (open access publication). *Genetics Selection Evolution*, 40(5) :1, 2008. [104](#)
- [24] J. A DiMasi, H. G Grabowski, and R. W Hansen. Innovation in the pharmaceutical industry : new estimates of r&d costs. *Journal of Health Economics*, 2016. [5](#)
- [25] J. A DiMasi, R. W Hansen, and H. G Grabowski. The price of innovation : new estimates of drug development costs. *Journal of health economics*, 22(2) :151–185, 2003. [5](#)
- [26] J. A DiMasi, R. W Hansen, H. G Grabowski, and L. Lasagna. Cost of innovation in the pharmaceutical industry. *Journal of health economics*, 10(2) :107–142, 1991. [5](#)
- [27] J. Donovan, M. R Elliott, and D. Heitjan. Predicting event times in clinical trials when randomization is masked and blocked. *Clinical Trials*, 4(5) :481–490, 2007. [92](#)
- [28] W. Feller. *An introduction to probability theory and its applications : volume I*, volume 3. John Wiley & Sons London-New York-Sydney-Toronto, 1968. [21](#)
- [29] A. Finkler. *Modèle d'évolution avec dépendance au contexte et Corrections de statistiques d'adéquation en présence de zéros aléatoires par*. PhD thesis, Université de Strasbourg, 2010. [104](#)
- [30] ICH Harmonised Tripartite Guideline. E6 (r1) guideline for good clinical practice : Consolidated guidance, ich june 1996. [2](#)
- [31] A. Hallstrom and K. Davis. Imbalance in treatment assignments in stratified blocked randomization. *Controlled Clinical Trials*, 9(4) :375–382, 1988. [84](#)

-
- [32] S. L Hedden, R. F Woolson, and R. J Malcolm. Randomization in substance abuse clinical trials. *Substance Abuse Treatment, Prevention, and Policy*, 1(1) :1, 2006. [84](#)
 - [33] Young J. Interim recruitment goals in clinical trials. *Journal of Chronic Diseases*, 36(5) :379–389, 1983. [6](#)
 - [34] N. Johnson, A. Kemp, and S. Kotz. *Univariate discrete distributions*, volume 444. John Wiley & Sons, 2005. [85](#)
 - [35] N. Johnson, S. Kotz, and N. Balakrishnan. Continuous univariate distributions, vol. 2 of wiley series in probability and mathematical statistics : applied probability and statistics, 1995. [21](#)
 - [36] N. L. Johnson, S. Kotz, and N. Balakrishnan. *Continuous univariate distributions. Vol. 2.* Wiley Series in Probability and Mathematical Statistics : Applied Probability and Statistics. John Wiley & Sons Inc., New York, second edition, 1995. A Wiley-Interscience Publication. [18](#), [19](#)
 - [37] N. L. Johnson, S. Kotz, and N. Balakrishnan. *Discrete multivariate distributions*, volume 165. Wiley New York, 1997. [21](#)
 - [38] N L. Johnson, S. Kotz, and A. W. Kemp. *Univariate discrete distributions.* Wiley Series in Probability and Mathematical Statistics : Applied Probability and Statistics. John Wiley & Sons Inc., New York, second edition, 1992. A Wiley-Interscience Publication. [42](#)
 - [39] J. M Lachin. Properties of simple randomization in clinical trials. *Controlled clinical trials*, 9(4) :312–326, 1988. [84](#)
 - [40] J. M Lachin. Statistical properties of randomization in clinical trials. *Controlled Clinical Trials*, 9(4) :289–311, 1988. [84](#)
 - [41] F. Le Gland. Introduction au filtrage en temps discret. *Filtre de Kalman, Modeles de Markov cachés.* IRISA/INRIA, 2003, 2002. [104](#)
 - [42] R. MacKay A. Assessing the goodness-of-fit of hidden markov models. *Biometrics*, 60(2) :444–450, 2004. [104](#)
 - [43] J. P Matts and J. M Lachin. Properties of permuted-block randomization in clinical trials. *Controlled clinical trials*, 9(4) :327–344, 1988. [84](#)
 - [44] G. Mijoule. *Modélisation du processus d’inclusion de patients dans un essai clinique multicentrique.* PhD thesis, Université de Toulouse, Université Toulouse III-Paul Sabatier, 2013. [12](#), [20](#)
 - [45] G. Mijoule, N. Minois, V. V Anisimov, and N. Savy. Additive cost modelling in clinical trial. In *Topics in Statistical Simulation*, pages 373–381. Springer, 2014. [8](#), [23](#), [95](#)
 - [46] G. Mijoule, S. Savy, and N. Savy. Models for patients’ recruitment in clinical trials and sensitivity analysis. *Statistics in Medicine*, 31(16) :1655–1674, 2012. [7](#), [12](#), [14](#), [15](#), [23](#), [53](#), [81](#), [88](#), [103](#)
 - [47] N. Minois, V. Lauwers-Cances, S. Savy, M. Attal, S. Andrieu, V. Anisimov, and N. Savy. Using poisson- gamma model to evaluate the duration of recruitment process when historical trials are available. *Statistics in Medicine*, 2016. In revision. [67](#), [103](#)
 - [48] N. Minois, G. Mijoule, V. Lauwers-Cances, S. Savy, S. Andrieu, and N. Savy. Performances of poisson- gamma model for patients ? recruitment in clinical trials

- when there are pauses in recruitment or when the number of centres is small. In *Proceedings of the 8th International Workshop on Simulation*, Vienne, 2015. [23](#)
- [49] N. Minois, S. Savy, V. Lauwers-Cances, S. Andrieu, and N. Savy. Poisson-gamma model for patients' recruitment in clinical trials. investigations on boundaries of relevancy by simulation studies. *Statistics in Biopharmaceutical Research*, 2015. Submitted. [52](#)
- [50] D. Moher, S. Hopewell, K. F. Schulz, V. Montori, P. C. Gotzsche, P.J. Devereaux, D. Elbourne, M. Egger, and D. G. Altman. Consort 2010 explanation and elaboration : updated guidelines for reporting parallel group randomised trials. *Journal of Clinical Epidemiology*, In Press, Corrected Proof :-, 2010. [4](#)
- [51] OECD. Serie sur les principes de bonnes pratiques de laboratoire et verification du respect de ces principes numéro 1. 1997. [2](#)
- [52] S. J Pocock. *Clinical trials : a practical approach*. John Wiley & Sons, 2013. [84](#)
- [53] W. F Rosenberger and J. M Lachin. *Randomization in clinical trials : theory and practice*. John Wiley & Sons, 2015. [84](#)
- [54] G. Samorodnitsky. A class of shot noise models for financial applications. In *Athens Conference on Applied Probability and Time Series Analysis, Vol. I (1995)*, volume 114 of *Lecture Notes in Statist.*, pages 332–353. Springer, New York, 1996. [96](#)
- [55] N. Savy and J. Vives. Anticipative integrals with respect to a filtered Lévy process and Lévy-Itô decomposition. Submitted to Communication on stochastic Analysis, 2016. [96](#)
- [56] S. Senn. *Statistical Issues in Drug Development*. John Wiley & Sons, Chichester, 1997. [6](#), [10](#), [77](#), [84](#)
- [57] S. Senn. Some controversies in planning and analysing multi-centre trials. *Statistics in Medicine*, 17 :1753–1765, 1998. [6](#)
- [58] A. Spagnoli, R. Henderson, R. J Boys, and J. J Houwing-Duistermaat. A hidden markov model for informative dropout in longitudinal response data with crisis states. *Statistics & Probability Letters*, 81(7) :730–738, 2011. [104](#)
- [59] R. Sukkar, E. Katz, Y. Zhang, D. Raunig, and B. T Wyman. Disease progression modeling using hidden markov models. In *2012 Annual International Conference of the IEEE Engineering in Medicine and Biology Society*, pages 2845–2848. IEEE, 2012. [104](#)
- [60] A.P. Vlasto. Brevets et médicament en france. pourquoi l'application du droit des brevets au médicament est-elle autant critiquée? *Médecine & Droit*, 2007(82) :25–32, 2007. [5](#)
- [61] W.O. Williford, S.F. Bingham, D.G. Weiss, J.F. Collins, K.T. Rains, and W.F. Krol. The "constant intake rate" assumption in interim recruitment goal methodology for multicenter clinical trials. *Journal of chronic diseases*, 40(4) :297–307, 1987. [6](#)
- [62] B. S. Yandell. *Practical Data Analysis for Designed Experiments*. Chapman & Hall, 1997. [30](#)
- [63] G. Ying and D. Heitjan. Weibull prediction of event times in clinical trials. *Pharmaceutical statistics*, 7(2) :107–120, 2008. [92](#)
- [64] G. Ying, D. Heitjan, and T. Chen. Nonparametric prediction of event times in randomized clinical trials. *Clinical Trials*, 1(4) :352–361, 2004. [92](#)