

Méthodes de traitement de l'information illustrées par une application au programme OCDE BEMUSE

- Journées des thèses IRSN, 1-4 octobre 2007 -

Sébastien Destercke

2^{ème} année de thèse

Méthodes de synthèse de l'information probabiliste imprécise

2 novembre 2005

thèse financée par l'IRSN

Directeur de thèse : Didier Dubois

Université Toulouse 3

Tuteur IRSN : Eric Chojnacki

DPAM/SEMIC/LIMSI

1 Contexte et objectifs

Lorsqu'on manipule des modèles physiques, il arrive souvent que l'information disponible sur des variables ou des paramètres soit pauvre, conflictuelle ou peu fiable. C'est d'autant plus vrai dans le cas des analyses de risques effectuées dans le domaine nucléaire, où la quantité de données expérimentales est très faible. Les théories probabilistes imprécises (possibilités [11], fonctions de croyances [15] et probabilités imprécises [17] elles-mêmes) offrent des outils permettant de représenter et traiter les cas où l'information concernant une situation est parcellaire, conflictuelle ou imprécise. Par rapport à la théorie des probabilités classiques, elles permettent de représenter l'information disponible de manière plus fidèle, évitant l'introduction d'hypothèses non-vérifiées.

Néanmoins, la généralité de ces modèles implique souvent une mise en oeuvre plus complexe, et un coût de calcul d'autant plus élevé. Afin de pouvoir traiter l'incertitude de manière rigoureuse et avec efficacité, il est donc indispensable de chercher à développer des méthodes de traitement de l'information utilisant les probabilités imprécises et dont la complexité soit en adéquation avec les besoins et les ressources disponibles pour un problème donné : alors que pour des modèles simples, des méthodes complexes de traitement de l'information pourront être utilisées, les coûts associés à des codes de calcul complexes ayant un nombre élevé de variables d'entrée impliqueront l'utilisation de représentations plus simples et de méthodes efficaces (quitte à être d'abord très conservatif pour affiner ensuite les résultats si cela s'avère nécessaire).

La figure 1 résume les principaux traitements que peut subir l'information concernant des variables ou des modèles. Chaque cercle représente un ensemble d'opérations qui peuvent être appliquées à l'information. Les relations liant les différents cercles montrent qu'il est possible d'effectuer des chaînes de traitement de différents types, la résolution d'un problème faisant souvent appel à plusieurs types d'opérations. Un exemple courant est celui d'un modèle ayant plusieurs variables d'entrée mal connues : si, pour chaque variable, plusieurs sources donnent des informations la concernant, il faut modéliser, fusionner, synthétiser ces informations, afin de pouvoir construire un modèle d'incertitude sur cette variable. Cette incertitude pourra ensuite être propagée à travers le modèle, afin d'estimer l'incertitude sur la(les) sortie(s) du modèle, et de prendre une décision en fonction de notre connaissance sur la(les) sortie(s). De même, l'estimation d'un paramètre d'une loi de probabilité (inférence) à partir de données expérimentales permet de construire un modèle qui pourra ensuite être propagé. Au final, le traitement d'informations entourées d'incertitudes passe donc par :

- Le choix d'une théorie de l'incertain,
- L'identification des outils à mettre en oeuvre.

Ces choix peuvent difficilement être faits a priori et demandent le plus souvent une bonne compréhension du problème traité et des besoins à satisfaire. La plupart du temps, la résolution du problème requiert donc une collaboration étroite entre l'(les) analyste(s) et l'(les) expert(s) du domaine concerné par l'analyse.

Les travaux de thèse se concentrent sur trois des quatre "boîtes" illustrées dans la figure 1 :

- Fusion, analyse, évaluation, synthèse [7] : étude du problème de la fusion de sources dépendantes entre-elles [8], proposition d'une méthode de fusion/analyse basée sur une approche utilisée en logique [10], application des méthodes de base au cas BEMUSE (benchmark OCDE) [5],
- Propagation : propagation numérique par méthode Monte-Carlo [2], calcul de moyennes supérieures/inférieures par intégrales [16],
- Modélisation : Etude des relations entre modèles pratiques de probabilités imprécises [9, 4], extension de méthodes numériques de simulation de dépendances [6].

Même si certains aspects de la prise de décision sont parfois pris en compte (voir [2]), les travaux de thèse accomplis se situent en amont des problématiques de décision, inférence et prédiction dans leur sens le plus général.

Pour illustrer l'intérêt du travail théorique accompli, nous exposerons les résultats obtenus par l'application de méthodes simples d'évaluation et de synthèse de sources multiples aux données du benchmark OCDE BEMUSE [13], au cours duquel plusieurs instituts ont comparé leurs résultats d'analyses d'incertitudes réalisées sur des codes de calculs simulant un accident sur un réacteur nucléaire.

Nous introduisons brièvement la problématique du traitement de l'information, et plus particulièrement de la fusion de sources multiples, dans la section 2. La section 3 explique la méthodologie générale d'évaluation et de synthèse des sources d'informations et la détaille pour les théories des probabilités et des possibilités (celles utilisées dans l'application). Enfin, la section 4 montre et commente les résultats obtenus pour le cas BEMUSE.¹

¹Une étude plus détaillée est réalisée dans [5]. Tous les calculs permettant l'évaluation et la synthèse des sources à partir des données fournies ont été réalisés au moyen du logiciel SUNSET

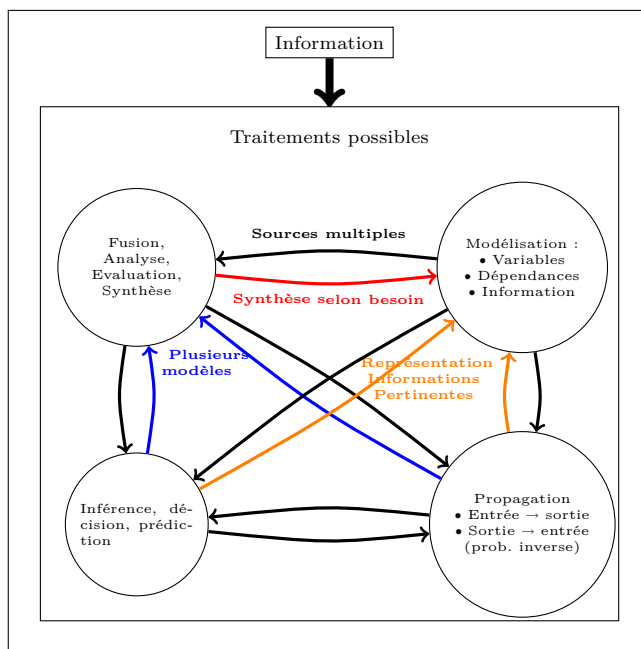


FIG. 1 – Méthodes de traitement de l'information - synoptique

2 Introduction à la problématique

Comme le montre la figure 1, le traitement de l'information implique de nombreux outils qui communiquent souvent entre-eux². En fonction de la situation (informations et ressources disponibles, demande du(des) décideur(s), ...), une partie ou l'intégralité de ces outils seront utilisés, à des niveaux de complexité divers. Alors que synthèse et propagation pourront être absentes du traitement, les étapes de modélisation et de décision seront toujours présentes : la modélisation est en effet indispensable à un traitement formel, et le but final de tout traitement est bien de prendre une décision.

Le fait de disposer d'un choix important d'outils de traitement permet de répondre à de nombreux problèmes, mais implique également un travail d'analyse important pour savoir quels outils sont adaptés à une situation donnée. Ainsi, l'étape de modélisation pourra être très simple si les modèles sont directement fournis ou si les données expérimentales sont nombreuses, mais pourra se révéler très complexe dans d'autres cas : élicitations et dialogues avec de nombreux experts qui peuvent être en conflit ; pauvreté des données disponibles ; information délivrée sous des formes hétérogènes (probabilités comparatives, intervalles de confiances, caractéristiques partielles d'une hypothétique mesure de probabilité, information qualitative ou linguistique, ...).

De même, la décision pourra se résumer à une réponse très simple (oui/non/besoin de plus d'informations) ou pourra consister en l'adoption d'une "politique" qui devra satisfaire des acteurs ayant des priorités différentes et qui sera constituée d'un ensemble d'options étalées dans le temps interagissant entre-elles. Une prédiction pourra être à très court terme et localisée (e.g. va-t-il pleuvoir cet après-midi ?) ou impliquera des processus macroscopiques et des enjeux à très long terme (e.g. évolution du climat). La propagation, quant à elle, pourra se faire à travers un modèle analytique simple, monotone en chaque variable ou à travers un code de calcul non-linéaire et coûteux (d'où le besoin de méthodes numériques de propagation qui cherchent à minimiser le nombre d'estimations à effectuer).

Dans la suite, nous nous concentrerons plus précisément sur la modélisation, la synthèse, l'évaluation et l'analyse de multiples sources d'information (la complexité de la synthèse dépendra du nombre de sources, du conflit entre-elles et de la complexité des modèles adoptés) : dans de nombreux domaines (et notamment le nucléaire), il arrive que plusieurs sources fournissent des informations concernant une même valeur mal connue (variable physique, paramètre, ...). C'est par exemple le cas quand plusieurs experts évaluent une variable physique mal connue, lorsque plusieurs senseurs fournissent des informations sur une même situation, ou encore lorsqu'il existe plusieurs modèles (codes de calcul, modèles analytiques) modélisant un même processus.

Lorsque cela arrive, évaluer la qualité des informations transmises, analyser ou synthétiser l'ensemble de ces informations de manière informelle sont souvent des tâches difficiles, voire impossibles. Le recours à des méthodes formelles de modélisation, évaluation, analyse et synthèse de l'information rend ces tâches plus aisées. De plus, leurs fondements théoriques permettent de justifier les résultats obtenus devant un décideur, ou de les communiquer plus facilement.

3 Méthodologies

Dans cette section, nous décrivons brièvement comment la méthodologie générale se décline dans les théories des probabilités et des possibilités (pour plus de détails, le lecteur est respectivement renvoyé à [1] et à [14])

²Le travail de thèse a notamment permis de clarifier les outils et les relations montrés par la figure

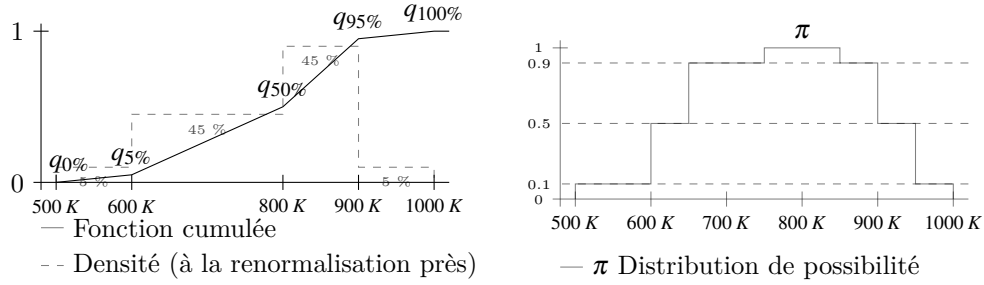


FIG. 2 – Modélisation : exemple

3.1 Modélisation

Lorsque les informations sont transmises sous forme probabiliste (intervalles de confiances, paramètres d'une distribution, percentiles, ...) et correspondent non pas à une mais à un ensemble de distributions possibles, il est courant d'utiliser le principe du maximum d'entropie. Ce principe permet d'obtenir une distribution unique correspondant à l'information partielle disponible tout en rajoutant un minimum d'information. Si ce choix est "optimal" lorsque l'on se contraint à rester au sein de la théorie des probabilités classiques, il ajoute néanmoins de l'information (non étayée empiriquement) à celle disponible et peut conduire à une dangereuse sous-estimation du risque. La théorie des possibilités offre alors une alternative simple et facile à interpréter et manipuler (son pouvoir expressif reste cependant limité, et il se peut que le recours à des modèles plus complexes de représentation de la connaissance soit nécessaire).

La figure 2 illustre les deux types de modélisations. La distribution de probabilité illustre le cas où une source fournit les valeurs 500 K, 600 K, 800 K, 900 K, 1000 K pour les percentiles à 0%, 5%, 50%, 95%, 100%, (la distribution de la figure 2 est celle correspondant à l'information donnée et ayant l'entropie maximum parmi l'infinité de distributions possibles). La distribution de possibilité représente le cas où la source fournit les intervalles [750 K, 850 K], [650 K, 900 K], [600 K, 950 K], [500 K, 1000 K] avec des niveaux de confiance respectifs de 10%, 50%, 90% et 100% (la distribution de possibilité correspondante englobe alors toutes les distributions de probabilités cohérentes avec l'information fournie).

3.2 Evaluation

L'évaluation d'une source se fait selon deux critères, appelés l'**informativité** et la **calibration**. L'informativité mesure la précision (et donc l'utilité potentielle) de l'information apportée par la source, par rapport au modèle non-informatif. La calibration se base sur des variables, appelées témoins, pour lesquelles les sources ont fourni des informations et dont les valeurs ont ensuite été observées (suite à des expériences) : la calibration consiste à mesurer la cohérence entre ces valeurs maintenant connues et l'information délivrée par la source. Une "bonne" source est donc une source d'information qui parvient à trouver un bon équilibre entre informativité et calibration, c'est-à-dire entre précision et adéquation avec l'expérience.

Théorie des probabilités : pour les probabilités, l'entropie est une mesure classique de disparité de l'information. C'est pourquoi on utilise une mesure dérivée de cette dernière pour mesurer l'informativité et la calibration d'une source. Soit $p = (p_1, \dots, p_B)$ et $q = (q_1, \dots, q_B)$ deux distributions de probabilités discrètes. La formule $I(p, q) = \sum_{i=1}^B p_i \log(p_i/q_i)$, appelée divergence de Kullbach-Leibler, permet de mesurer la différence de p par rapport à q . Elle est utilisée à la fois pour mesurer l'informativité et la calibration d'une source. Pour l'informativité, p est la distribution correspondant aux informations délivrées par la source, et q est la distribution uniforme (modèle représentant l'ignorance dans la théorie des probabilités). Pour la calibration, p est une distribution empirique construite à partir de valeurs observées, et q est la distribution correspondant aux informations délivrées par la source. Il est intéressant de noter que comparer deux probabilités requiert une série de données, ce qui impose une modélisation commune (i.e. la même série de percentiles) sur chaque variable témoin. Cette limitation n'est pas nécessaire dans les autres théories de l'incertain.

Théorie des possibilités : soit π la distribution de possibilité construite à partir des informations données par la source, L, U les bornes inférieures et supérieures des valeurs que peut potentiellement prendre la variable (avant de connaître sa vraie valeur), et x^* la valeur expérimentale prise par cette variable. L'informativité se calcule par la différence $(U - L) - \left(\int_L^U \pi(x) dx\right)$, tandis que la calibration est mesurée par la valeur $\pi(x^*)$ (qui représente la plausibilité attribuée à la valeur x^*).

L'ensemble des scores de calibration et informativité sont ensuite agrégés pour attribuer un poids w_i à la $i^{\text{ème}}$ source, tels que $\sum_{i=1}^n w_i = 1$, avec n le nombre total de sources (i.e. la somme des poids est normalisée).

3.3 Synthèse

Une fois les poids calculés, il existe trois opérations basiques permettant de synthétiser l'information délivrée par les différentes sources (synthèse pouvant être utilisée pour analyser l'information ou simplement pour construire un modèle "résumant" l'information disponible) :

- la conjonction : équivalent de l'intersection, elle donne des résultats informatifs, mais peu fiables en cas de conflits,
- la disjonction : équivalent de l'union, elle donne des résultats très fiables mais généralement trop imprécis pour être réellement utiles,
- la moyenne pondérée : elle revient à faire un comptage entre les sources en supposant ces dernières indépendantes entre elles (hypothèse non vérifiée la plupart du temps).

Pour les probabilités, la seule opération possible est la moyenne pondérée (les deux autres opérations étant des opérations ensemblistes) : elle revient à calculer $p_{\text{moy}} = \sum_{i=1}^n w_i p_i$, avec p_i la distribution correspondant à la $i^{\text{ème}}$ source.

Dans le cadre des possibilités, les trois opérations sont possibles : la conjonction, la disjonction et la moyenne pondérée reviennent respectivement à calculer $\pi_{\cap}(x) = \min_{i=1, \dots, n} (\pi_i(x))$, $\pi_{\cup}(x) = \max_{i=1, \dots, n} (\pi_i(x))$ et $\pi_{\text{moy}}(x) = \sum_{i=1}^n w_i (\pi_i(x))$. L'approche possibiliste (et, en général, les approches ensemblistes adoptées par les autres théories de l'incertain) permet donc une plus grande flexibilité dans la synthèse et l'analyse de l'information. La figure 3 illustre les synthèses possibles pour les deux théories.

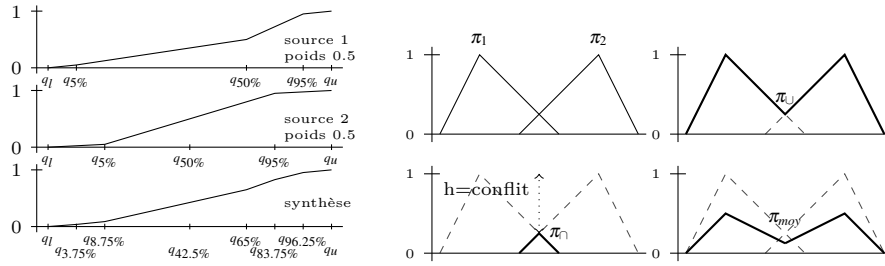


FIG. 3 – Synthèse : exemple

4 Application

BEMUSE est un programme OCDE dont l'objectif est d'évaluer les méthodes d'analyse d'incertitudes sur un accident de dimensionnement (licensing). Lors de ce programme, 9 instituts ont dû estimer, grâce à leurs codes de calcul respectifs, l'incertitude sur les valeurs prises par plusieurs variables dans l'hypothèse d'un accident de perte de refroidissement consécutif à une grosse brèche. Chaque institut peut être considéré comme une source d'information, et ses résultats sur chaque variable comme les informations transmises par cette source. Sous cet éclairage, l'ensemble des informations et résultats peuvent donc être analysés et traités par les méthodes présentées dans la section précédente.

Les variables que nous avons retenues sont les deux valeurs de pic de température dans la gaine de crayons combustibles (PCT dans le tableau 1), le temps d'injection d'eau de secours (T_{inj}) et le temps de renoyage du coeur (T_q). Les valeurs expérimentales ont été obtenues sur la boucle LOFT. Le tableau 1 résume les résultats obtenus et communiqués par les différents instituts. Pour chaque variable, une borne inférieure (Low), une valeur de référence (Ref) et une borne supérieure (Up) ont été fournies. Nous avons interprété les valeurs respectivement comme les percentiles 2, 50 et 98 %, et leur avons associé les modèles probabilistes et possibilistes correspondants. Les valeurs expérimentales étant disponibles, nous avons pu procéder à l'évaluation de la qualité des informations transmises.

4.1 Evaluation

Le tableau 2 résume les résultats obtenus lors des procédures d'évaluation, à la fois pour l'approche probabiliste et l'approche possibiliste.

Du point de vue des méthodes, nous constatons une bonne cohérence entre les deux classements (parmi les cinq premiers des deux classements, quatre sont communs aux deux : IRSN, KINS, NRI1, UNIP1), ce qui est logique, vu que les deux approches suivent les mêmes principes conceptuels (c'est particulièrement vrai pour la mesure d'informativité). Il existe cependant des différences mises en évidence par l'application, et qui sont dues aux différences de mises en oeuvre de la méthodologie dans chaque théorie. Ces différences sont notamment dues :

- au fait que les dyssymétries jouent un rôle important dans le calcul de l'informativité probabiliste, et moins important pour l'approche possibiliste (d'où les résultats différents obtenus pour KINS).
- au fait que la calibration probabiliste se base sur un argument de convergence vers une distribution "idéale", tandis que la calibration possibiliste se base sur la distance de la valeur expérimentale par rapport à la valeur de référence (d'où les différences de classement notables pour GRS et NRI2, selon l'approche adoptée)

Du point de vue de l'analyse d'information, on remarque que les instituts ne sont pas groupés par code utilisé (l'IRSN et le CEA, utilisateurs de CATHARE, ont des scores assez différents, ainsi que les utilisateurs des codes ATHLET et RELAP5 entre eux). Ce résultat confirme que, plus que le code utilisé, c'est la façon de l'utiliser qui importe. Les résultats obtenus sont également cohérents avec les observations informelles, vu qu'UPC et PSI, les deux seuls instituts pour qui les valeurs expérimentales sont en dehors de l'intervalle $[Low, Up]$ pour deux variables, se retrouvent dans le bas du classement.

	Code	1PCT (K)			2PCT (K)			T_{inj} (s)			T_q (s)		
		Low	Ref	Up	Low	Ref	Up	Low	Ref	Up	Low	Ref	Up
CEA	CATHARE	919	1107	1255	674	993	1176	14.8	16.2	16.8	30	69.7	98
GRS	ATHLET	969	1058	1107	955	1143	1171	14	15.6	17.6	62.9	80.5	103.3
IRSN	CATHARE	872	1069	1233	805	1014	1152	15.8	16.8	17.3	41.9	50	120
KAERI	MARS	759	1040	1217	598	1024	1197	12.7	13.5	16.6	60.9	73.2	100
KINS	RELAP5	626	1063	1097	608	1068	1108	13.1	13.8	13.8	47.7	66.9	100
NRI1	RELAP5	913	1058	1208	845	1012	1167	13.7	14.7	17.7	51.5	66.9	87.5
NRI2	ATHLET	903	1041	1165	628	970	1177	12.8	15.3	17.8	47.4	62.7	82.6
PSI	TRACE	961	1026	1100	887	972	1014	15.2	15.6	16.2	55.1	78.5	88.4
UNIP1	RELAP5	992	1099	1197	708	944	1118	8.0	16.0	23.5	41.4	62.0	81.5
UPC	RELAP5	1103	1177	1249	989	1157	1222	12	13.5	16.5	56.5	63.5	66.5
Exp. Val.		1062			1077			16.8			64.9		

TAB. 1 – Estimations des valeurs de sortie par participants (Exp. Val. : valeur expérimentale)

	Approche prob.			Approche Poss.		
	Inf.	Cal.	Global	Inf.	Cal.	Global
CEA	8 (0.77)	5 (0.16)	6 (0.12)	8 (0.71)	6 (0.55)	7 (0.40)
GRS	4 (1.23)	1 (0.98)	1 (1.21)	3 (0.84)	7 (0.52)	6 (0.44)
IRSN	5 (0.98)	2 (0.75)	2 (0.73)	6 (0.73)	1 (0.83)	1 (0.60)
KAERI	9 (0.68)	5 (0.16)	7 (0.11)	9 (0.70)	8 (0.48)	8 (0.34)
KINS	3 (1.29)	5 (0.16)	5 (0.21)	7 (0.72)	3 (0.67)	3 (0.49)
NRI1	7 (0.79)	2 (0.75)	3 (0.59)	5 (0.75)	5 (0.63)	4 (0.47)
NRI2	6 (0.79)	8 (0.13)	8 (0.10)	4 (0.78)	2 (0.72)	2 (0.56)
PSI	1 (1.6)	10 (0.004)	10 (0.008)	1 (0.88)	10 (0.25)	10 (0.22)
UNIP1	10 (0.53)	2 (0.75)	4 (0.4)	10 (0.69)	4 (0.67)	5 (0.46)
UPC	2 (1.44)	9 (0.02)	9 (0.025)	2 (0.87)	9 (0.28)	9 (0.24)

TAB. 2 – Evaluation des sources (Inf. : informativité ; Cal. : calibration) par rangs (valeurs)

4.2 Synthèse

Les figures 4.A-F résument les résultats obtenus pour la synthèse concernant le second pic de température (2PCT) qui est la variable la plus intéressante à analyser, à la fois parce qu'elle est la plus difficile à estimer et la plus importante en terme de sûreté.

Les résultats de l'approche probabiliste (fig 4.A) indiquent que les utilisateurs de RELAP5 et de CATHARE tendent à sous-estimer le second pic, alors que les utilisateurs ATHLET ont plutôt tendance à le surestimer. La figure 4.B montre l'effet de la pondération ainsi que celui de se restreindre aux sources les mieux classées : à la fois les quatre premières de l'approche probabiliste (GRS, NRI1, UNIP1, IRSN) et les meilleures sources communes aux deux classements (IRSN, KINS, NRI1, UNIP1). On observe un léger resserrement des distributions pondérées (gain d'informativité) ainsi qu'un décalage vers la valeur expérimentale (gain de calibration). Cependant, le peu de différence entre les différentes distributions montre que l'utilisation de la moyenne tend à lisser le résultat (voir aussi la figure 4.C), ce qui montre que cet opérateur de synthèse est peu discriminant. De plus, le conflit entre sources n'est pas visible par cette approche.

Grâce au plus grand nombre d'opérateurs de synthèse disponibles, l'approche possibiliste permet de mettre en évidence d'autres informations. La disjonction (fig 4.B) confirme que les sources tendent à sous-estimer ou à surestimer la température de pic. Le "trou" observable autour de la valeur expérimentale indique que la source la plus proche de la valeur expérimentale (KINS) est également très (trop ?) précise. Comme prévu, la distribution 4.B est très imprécise, mais juge la valeur expérimentale comme plausible à 80 %. Les résultats obtenus en utilisant la conjonction (fig 4.E-F) permettent quand à eux de mettre en évidence les conflits existants entre les sources. Ainsi, les utilisateurs de CATHARE montrent une grande cohérence dans leurs résultats (niveau de conflit inférieur à 0.1), cohérence qui est moindre pour les utilisateurs d'ATHLET (niveau de conflit autour de 0.4) et de RELAP5 (le plus grand conflit pouvant s'expliquer par le plus grand nombre de sources dans ce dernier cas). La figure 4.F montre que le conflit sur l'ensemble des sources est très élevé (~0.9!), par contre, les sources les mieux classées montrent une grande cohérence, ce qui, d'un point de vue pratique, nous conforte dans le choix de ces sources.

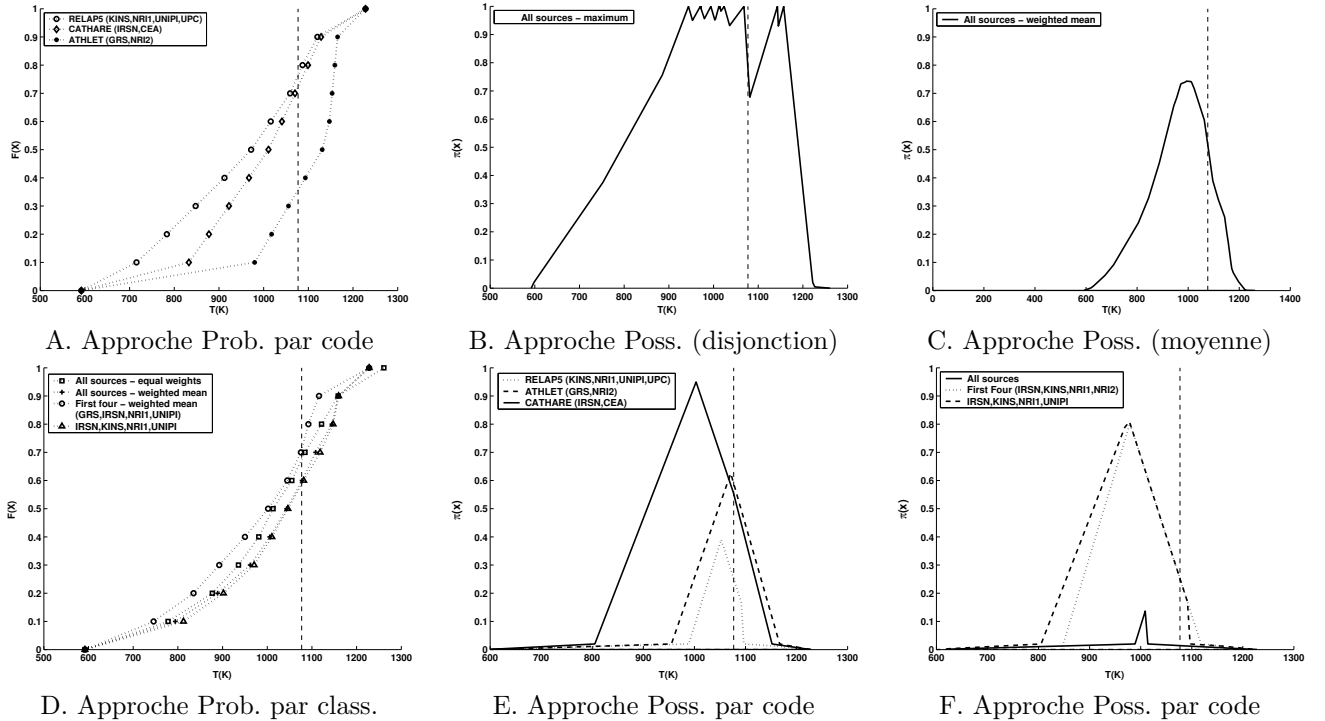


FIG. 4 – Résultats de la synthèse : selon codes et classements du tableau 2 (--- : val. exp.)

5 Conclusions et perspectives

Nous avons appliqué les méthodes usuelles d'évaluation et de synthèse de sources multiples d'information au benchmark BEMUSE. D'un point de vue méthodologique, les points communs et les divergences entre les approches probabiliste et possibiliste ont été mis en évidence, notamment au travers de l'application.

D'un point de vue pratique, on a observé la concordance entre méthodes formelles et observations informelles, ce qui montre que les formalismes sont bien adaptés au traitement du problème. L'utilisation des méthodes formelles a permis de mettre en évidence des informations difficiles à extraire des données brutes, tels les conflits existants entre sources, ou encore la tendance à sous-estimer, surestimer certaines valeurs.

Nous avons montré que même des méthodes simples de traitement de l'information présentent un intérêt et peuvent apporter un nouvel éclairage sur certains problèmes. C'est pourquoi il est utile de développer des méthodes formelles plus fines, basées sur des arguments à la fois séduisants intuitivement et mathématiquement justifiés, qui permettent de réaliser des études plus complètes ou encore de traiter des situations plus complexes. La méthode de synthèse basée sur les sous-ensembles maximaux cohérents, actuellement étudiée et développée dans le cadre de la thèse (voir [10]), répond (au moins partiellement) à ce besoin.

Au cours de cette seconde année de thèse, 8 articles ont été acceptés à diverses conférences, certains étant actuellement retravaillés pour être soumis à des journaux. Outre l'écriture du mémoire, la troisième année de thèse sera vraisemblablement consacrée à une étude plus approfondie des diverses notions d'indépendance [3, 12, 18] existant dans les diverses théories de l'incertain, tout en approfondissant les aspects du traitement de l'information déjà abordés. Ces efforts pourront se concrétiser par des publications (conférences et versions longues d'articles déjà publiés) et/ou l'implémentation de nouvelles méthodes dans SUNSET. La formation par la recherche sera de plus complétée par une deuxième participation à un comité d'organisation de conférence (SMPS'08), ainsi que par la participation à un projet de livre portant sur les méthodes récentes de traitement de l'incertitude.

Références

- [1] T. Bedford and R. Cooke. *Probabilistic Risk Analysis. Foundations and Methods*. Cambridge University Press, UK, 2001.
- [2] E. Chojnacki, J. Baccou, and S. Destercke. Numerical sensitivity and efficiency in the treatment of epistemic and aleatory uncertainty. In *Proc. Fifth Int. Conf. on Sensitivity Analysis of Model Output*, 2007.
- [3] I. Couso, S. Moral, and P. Walley. A survey of concepts of independence for imprecise probabilities. *Risk Decision and Policy*, 5 :165–181, 2000.
- [4] G. de Cooman, E. Miranda, M. Troffaes, S. Destercke, and D. Dubois. Notes on probability boxes and related models. *In progress*, 2007.
- [5] S. Destercke and E. Chojnacki. Evaluation and synthesis of multiple sources of information : an application to nuclear computer codes. *Submitted to Nuclear Eng. and Design*, 2007.
- [6] S. Destercke and E. Chojnacki. Extending usual methods handling dependencies to imprecise probabilistic models. In *Submitted to ninth Int. probabilistic safety assessment and management conf.*, 2008.
- [7] S. Destercke, D. Dubois, and E. Chojnacki. Fusion d'opinions d'experts et théories de l'incertain. In *Proc. Rencontres Francophones sur la logique floue et ses applications*, 2006.
- [8] S. Destercke, D. Dubois, and E. Chojnacki. Cautious conjunctive merging of belief functions. In *Proc. Eur. Conf. on Symbolic and Quantitative Approaches to Reasoning with Uncertainty*, 2007.
- [9] S. Destercke, D. Dubois, and E. Chojnacki. On the relations between practical representations of imprecise probabilities. *Submitted to Int. J. of Approximate Reasoning*, 2007.
- [10] S. Destercke, D. Dubois, and E. Chojnacki. Possibilistic information fusion using maximal coherent subsets. In *Proc. IEEE Int. Conf. On Fuzzy Systems (FUZZ'IEEE)*, 2007.
- [11] D. Dubois, H. Nguyen, and H. Prade. *Fundamentals of Fuzzy Sets*, chapter Possibility theory, probability and fuzzy sets : misunderstandings, bridges and gaps, pages 343–438. D. Dubois and H. Prade, 2000.
- [12] T. Fetz and M. Oberguggenberger. Propagation of uncertainty through multivariate functions in the framework of sets of probability measures. *Reliability Engineering and System Safety*, 85 :73–87, 2004.
- [13] OCDE. Bemuse phase iii report : Uncertainty and sensitivity analysis of the loft 12-5 test. Technical Report NEA/NCIS/R(2007)4, NEA, May 2007.
- [14] S. Sandri, D. Dubois, and H. Kalfsbeek. Elicitation, assessment and pooling of expert judgments using possibility theory. *IEEE Trans. on Fuzzy Systems*, 3(3) :313–335, August 1995.
- [15] G. Shafer. *A mathematical Theory of Evidence*. Princeton University Press, 1976.
- [16] L. Utkin and S. Destercke. Computing expectations with p-boxes : two views of the same problem. In *Proc. of the fifth Int. Symp. on Imprecise Probabilities and Their Applications*, 2007.
- [17] P. Walley. *Statistical reasoning with imprecise Probabilities*. Chapman and Hall, 1991.
- [18] B. B. Yaghlane, P. Smets, and K. Mellouli. Belief function independence : I. the marginal case. *I. J. of Approximate Reasoning*, 29(1) :47–70, 2002.