

ALASCA — Architecture logicielles avancées
pour les systèmes cyber-physiques autonomiques

© Jacques Malenfant

Master informatique, spécialité STL – UFR 919 Ingénierie

Sorbonne Université
Jacques.Malenfant@lip6.fr

Cours 10

**De l'entité de contrôle —
méthodologie et développement**

Objectifs pédagogiques du cours 10

- Au-delà des notions plus conceptuelles vues au cours précédent, acquérir des savoir-faire pratiques pour la mise en œuvre des entités de contrôle.
- Apprendre comment programmer les contrôles et les différentes fonctions de l'entité de contrôle.
- Introduire des approches permettant de rendre la composition entre entités adaptables et entités de contrôle plus robustes.

Construction des entités de contrôle

- Dans ce cours, nous allons nous intéresser aux entités de contrôle *individuelles* implantant un contrôle *unique* ou la *fusion* de plusieurs contrôles en *un seul modèle* de contrôle.
- La construction des contrôleurs *composites* ou *coordonnés* sera essentiellement étudiée au cours 12, mais nous introduirons ici des éléments utiles à la conception des contrôleurs *individuels* en prévision de leur composition et de leur coordination.
- Bien que les entités de contrôle puissent être implantées de manière réactive ou délibérative, nous allons reprendre un examen structuré autour des quatre grandes fonctions des gestionnaires autonomiques d'IBM : supervision, analyse/décision, planification et exécution des adaptations.
- Nous traiterons enfin de problématiques au niveau du processus de développement et du cycle de vie : test, validation, auto-configuration.

Langages spécialisés

- La programmation impérative des contrôles multiples suppose d'organiser les calculs afférents aux différents contrôles dans un même cycle de contrôle.
- Les contrôleurs réactifs sont souvent programmés à partir de blocs de code (en anglais *code elements* ou *codels*) devant être ordonnancés en fonction :
 - de leurs contraintes temporelles (délais de réaction, ...),
 - de leurs relations de précédences (production/consommation de données, ...)
 - et d'autres contraintes d'accès aux ressources.
- Quelques langages spécialisés ont été proposés pour faciliter l'expression des relations utilisées pour l'ordonnancement, que nous allons illustrer par deux exemples :
 - les langages synchrones en programmation temps-réel, et
 - les architectures comportementales de Brooks en robotique.

Les langages synchrones I

- Les langages synchrones ont été proposés pour programmer des systèmes réactifs, c'est-à-dire réagissant en permanence avec un délai borné à des événements se produisant dans leur environnement par des actions sur ce dernier.
- Les langages synchrones sont fondés sur un ensemble commun d'hypothèses dont les principales sont :
 - l'existence d'un temps discret où tous les événements se produisent, les réactions calculées et exécutées ;
 - les calculs et les réactions sont instantanés, c'est-à-dire ont une durée d'exécution nulle.
- Il existe deux familles (« duales ») de langages synchrones :
 - ① les langages à flût de données, et
 - ② les langages dirigés par les événements, que nous allons plus particulièrement survoler.

Les événements peuvent être des données reçues des capteurs.

Les langages synchrones II

- Par leurs hypothèses, les langages dirigés par les événements :
 - supposent des interventions à des instants discrets où l'ensemble des événements qui se sont produits depuis la dernière intervention sont « consommés » ;
 - le calcul des réactions (*i.e.*, actions d'adaptation) s'assimile alors à l'exécution d'un automate où les événements provoquent des transitions et l'exécution de *code/s* émettant de nouveaux événements ou appelant des actionneurs.
- Par rapport à une expression directe sous la forme d'automates, les langages synchrones facilitent la production de ces derniers par une expression supposée plus facile à écrire, qui sera ensuite compilée vers l'automate correspondant.

L'automate est un peu un langage intermédiaire du compilateur de langage synchrone...

Exemple : compteur de vitesse

```

module Speedometer:
  input Second, Meter ; output Speed
  in loop
    var Distance:= 0: integer
    in do every Meter do
      Distance:= Distance + 1
    end every
    upto Second ;
    emit Speed(Distance)
  end var
end loop
end module.

```

```

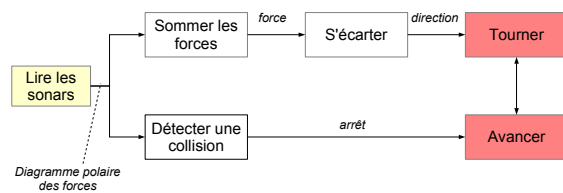
module SpeedSupervisor
input Speed, Meter ; output TooFast
in signal Speed: integer in [
    run Speedometer
    ||
    every Speed do
        if ?Speed > MaxSpeed then
            emit TooFast
        end if
    end every
]
end signal
end module

```

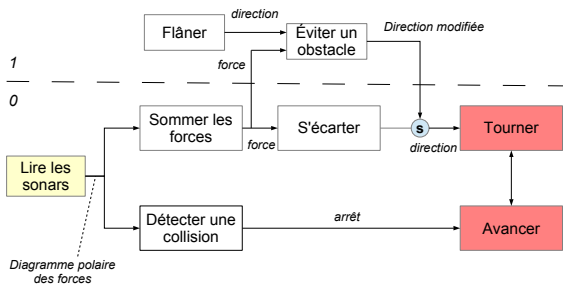
- Parmi les nombreuses constructions non-illustrées par cet exemple, il est possible d'appeler des *code/s* externes.
- L'opérateur de parallélisme '||' sert à exprimer un parallélisme *logique*, permettant de composer plusieurs contrôles sans les ordonner explicitement et sans qu'il y ait exécution parallèle.

Exemple : un robot suivant un corridor I

- 1 Niveau 0 : avancer sans entrer en collision avec les obstacles.

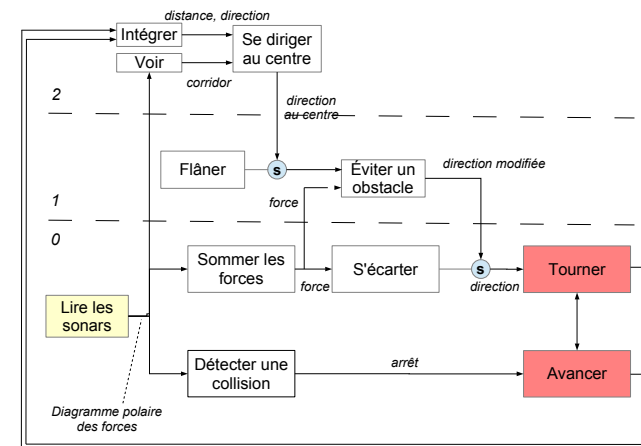


- 2 Niveau 1 : flâner, tout en évitant les obstacles.



Exemple : un robot suivant un corridor II

- 1 Niveau 3 : suivre un corridor en évitant les obstacles.



Quelques observations sur les architectures à subsomption

- Par la synchronisation des lectures de capteurs en début de cycle, l'architecture à subsomption résout le problème du partage des capteurs entre différents contrôles.
- Par les relations de précédences, cette architecture permet l'ordonnancement des tâches à l'intérieur du cycle et la vérification de son ordonnancement globale.
- Par son organisation en couches et les mécanismes d'inhibition et de suppression, elle compose les comportements avec des priorités explicites, mais cette technique est purement *ad hoc* et donc très difficile de prouver des propriétés sur le contrôleur.
- Par la composition des signaux d'actions, comme les champs de potentiels, cette architecture (comme d'autres) autorise aussi le déclenchement de plusieurs actions se combinant en fin de cycle avant l'envoi des commandes aux actionneurs.
Ex.: des directions calculées comme des vecteurs sont simplement additionnées pour produire la nouvelle direction.

Représenter γ : approches fonctionnelles et statistiques

- Dans les exemples précédents, la loi de commande a été exprimée par du code or dans beaucoup de cas, la loi est trop complexe pour être facilement programmée de cette manière.
- Une loi de commande $\gamma : \mathbf{S} \rightarrow \mathbf{A}$ peut être connue par un nuage de points, suite à des expérimentations et des mesures à l'exécution, par exemple.
- Si γ est continue, trouver une représentation à partir de points $\{(\mathbf{s}_i, \mathbf{a}_i)\}$ est une problématique étudiée en analyse numérique.
- Types de représentations :
 - paramétriques : fonctions approximées (linéaires, splines, séries de Taylor, séries de Fourier, ...) possiblement décomposées en morceaux (linéaires ou splines par morceaux, ...);
 - non-paramétriques : tableaux, structures de données associatives (tables de hachage), diagrammes de décision binaires, règles floues, réseaux de neurones, ...

Approximations paramétriques

- Il s'agit d'utiliser une forme de fonction prédéfinie avec des paramètres $p_0, \dots, p_k$ puis à trouver *globalement* les valeurs de ces paramètres qui collent (le mieux) avec l'ensemble des points.
- Comme le problème devient vite surcontraint, il faut généralement utiliser des approximations qui ne passent pas nécessairement par tous les points.
- Les approximations *par morceaux* permettent d'améliorer les choses en trouvant les meilleures valeurs *localement à chaque morceau*, n'utilisant qu'un nombre limité de points à chaque fois. Ex.: une fonction linéaire par morceau définie sur les points (x_i, y_i) en ordre croissant de x_i :

$$y = \frac{(y_{i+1} - y_i)}{(x_{i+1} - x_i)}x + \frac{(x_{i+1}y_i - x_iy_{i+1})}{(x_{i+1} - x_i)}, \quad x_i \leq x \leq x_{i+1}$$

Approximations non paramétriques

- Les représentations non-paramétriques les plus simples sont les tableaux ou les tables de hachage qui fonctionnent plutôt dans le cas des *fonctions discrètes* où *tous les points sont connus*.
- Dans le cas de *fonctions continues*, une *interpolation* deviendra nécessaire, comme dans les approximations paramétriques. Comme dans ce cas, on peut se fonder sur les intervalles et utiliser des diagrammes de décision binaires ou toutes autres structures arborescentes pour accéder plus rapidement à l'intervalle contenant le s pour lequel la valeur de a est requise.
- Ces approches simples deviennent cependant trop coûteuses en mémoire lorsque le nombre de points devient trop grand.
- Pour réduire la taille mémoire nécessaire, on utilise des représentations capables de *minimiser* la quantité d'informations nécessaire, dont par exemple les *réseaux de neurones profonds* qui connaissent une forte popularité.

Approches statistiques

- Les systèmes réels sont rarement déterministes :
 - des erreurs ou des dérives aléatoires peuvent entacher la trajectoire d'évolution de l'état ;
 - les conséquences des actions peuvent être entachées d'aléa ;
 - les meilleures lois de commande ou politiques peuvent consister à prendre des décisions aléatoires.
- Pour les deux premiers cas, il faut prendre en compte de l'aléa dans le modèle d'évolution de l'état du système.
 - Par exemple, en étudiant les contrôleurs PID nous avons supposé que l'évolution du système pouvait être modélisée par une fonction $p(k+1) = a_0 p(k) + a_1 u(k)$.
 - Comment déterminer les paramètres comme a_0, a_1 ?
 - Des expérimentations peuvent fournir des nuages de points.
 - Il faut ensuite estimer les paramètres en fonction de ces points.
- Les réseaux de neurones peuvent représenter des modèles stochastiques, mais nous allons maintenant plutôt voir des méthodes purement statistiques.

Modèles de régression

- Des modèles relativement simples qui s'appuient sur la notion de *corrél*ation entre variables aléatoires :
 - une corrélation existe entre variables aléatoires observées sur une même population lorsque les variations de ces variables se produisent dans le même sens ou en sens contraire.
- *I*dée : considérer le contrôle U et l'état du système P comme deux variables aléatoires corrélées.
- Le *coefficient de corrél*ation linéaire est une mesure de l'*intensité de la liaison* qui existe entre ces variables aléatoires. Soient les couples (x_i, y_i) , on calcule le coefficient r par la formule :

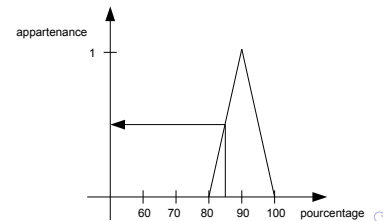
$$r = \frac{\sum_{i=1}^n (x_i - \bar{x})(y_i - \bar{y})}{\sqrt{\sum_{i=1}^n (x_i - \bar{x})^2} \sqrt{\sum_{i=1}^n (y_i - \bar{y})^2}} \quad \begin{aligned} \bar{x} &= \frac{\sum_{i=1}^n x_i}{n} \\ \bar{y} &= \frac{\sum_{i=1}^n y_i}{n} \end{aligned}$$

- r varie dans l'intervalle $[-1, 1]$, et plus $|r|$ est proche de 1, plus la corrélation est forte.

Introduction (très rapide) à la logique floue

- L'objectif de la logique floue est de capturer l'imprécision des données et d'assurer un raisonnement dans l'imprécision *bien fondé*.
- Les domaines de valeurs sont subdivisés en sous-ensembles flous représentant les « qualités », comme un pourcentage considéré comme TRÈS-FAIBLE, FAIBLE, MOYEN, ÉLEVÉ, TRÈS-ÉLEVÉ.
- C'est une extension de la logique classique, où l'appartenance à un ensemble (et donc les valeurs de vérité) n'est plus binaire (vrai ou faux) mais un degré entre 0 et 1.
- Un sous-ensemble flou $S_i \in S$ est défini par le degré d'appartenance des valeurs de S à S_i , la multi-appartenance étant admise.

Ex. : Définissons ci-contre un sous-ensemble flou de pourcentage dit TRÈS-ÉLEVÉ. Le degré d'appartenance de la valeur 85% est de 0.5.



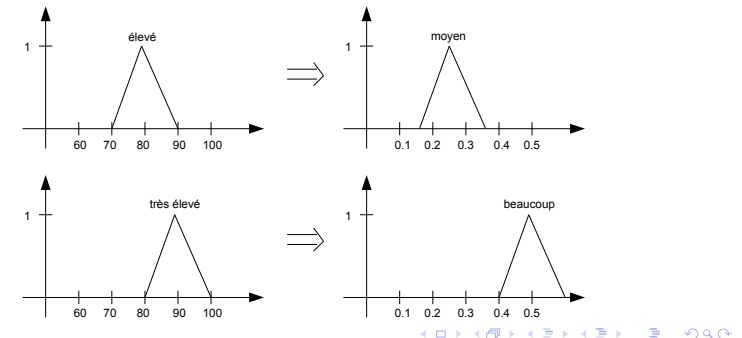
Règles d'inférence

- Une règle d'inférence floue a pour prémisses et conclusion des sous-ensembles flous.

Ex. : Définissons deux règles

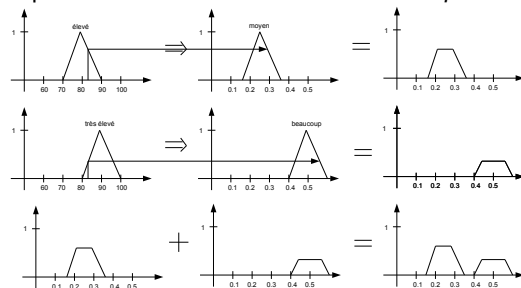
R1 : si (pourcentage) élevé alors (quantité de mémoire allouée) moyen

R2 : si très élevé alors beaucoup



Raisonnement flou

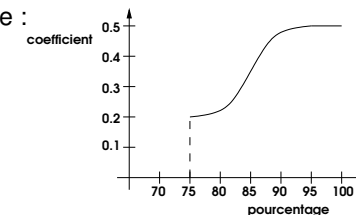
- Différentes procédures de raisonnement ont été proposées pour la logique floue, qui toutes ont pour objectif d'être compatibles avec la logique classique dans le cas *sans* imprécision, et de fournir des résultats satisfaisants dans le cas *avec* imprécision.
- Un exemple de raisonnement très visuel avec $p = 82.5\%$:



Puis, *defuzzification* par centre de gravité donne un coefficient d'environ 0.35.

Autres intérêts du contrôle flou

- Cette *discretisation souple* produit des résultats continus et progressifs. Exemple :



- La procédure d'inférence utilisant toutes les règles applicables pour fondre en un résultat unique tous les résultats des différentes règles fournit *de facto* un *mécanisme de combinaison* pouvant être pertinent dans certains cas de contrôles multiples.
- Comme déjà mentionné, l'approche floue est beaucoup utilisée en contrôle industriel et il existe des bibliothèques complètes dans de nombreux langages, dont Java.

- 1 Fonction de supervision
- 2 Fonction de contrôle et décision
- 3 Fonction de planification**
- 4 Fonction d'adaptation

Planification

Activité consistant à produire un enchaînement d'actions valides permettant d'atteindre un but à partir d'un état initial donné.

- Deux grandes approches de planification :
 - la *recherche dans un espace d'états* manipule un *graphe des états*, où les transitions sont des actions, pour explorer différents agencements temporels d'actions (séquences, alternatives, répétitions) et choisir celui la plus apte à passer de l'état initial à un état but ;
 - la *recherche dans l'espace des plans* manipule des *plans partiels puis complets*, sélectionnant ou remplaçant progressivement les actions les plus susceptibles de compléter ou d'améliorer le plan partiel pour se rapprocher du plan final ou optimal.

- Quand le plan est sous-contraint, les marges de manœuvre dans les actions sont préservées en produisant les plans les plus *permissifs* possibles.
- Le planificateur comporte :
 - un *modèle déclaratif de planification* décrivant le système, son espace d'états, ses actions et leurs contraintes,
 - un *moteur de planification* raisonnant sur le modèle de planification pour produire des plans,
 - Un *moteur d'exécution* des plans ordonnant, exécutant et supervisant la réalisation des actions.

Idéalement, modèle et moteurs sont indépendants, mais en pratique ils sont souvent liés par des heuristiques et techniques de raisonnement rendant les planificateurs de différents modèles incompatibles les uns avec les autres.

- La supervision de l'exécution des plans peut constater l'échec d'un plan, auquel cas il est possible de répondre par *replanification* ou par *réparation* de plan puis poursuite ou ré-exécution du plan modifié.

Plans et planificateurs

- La planification est l'un des éléments-clés de l'intelligence artificielle *située*, c'est-à-dire la conception d'agents (logiciels ou physiques) capables d'agir de manière intelligente et autonome dans un environnement, comme les robots autonomes.
- Malheureusement, la complexité exponentielle de la planification en général n'a pas permis de produire un planificateur *universel*, utilisable pour *tous* les problèmes.
- La planification a donc fait l'objet de très nombreuses recherches, utilisant différents formalismes pour différentes classes de problèmes, avec des résultats théoriques et pratiques utilisables dans le domaine des systèmes auto-adaptables.
- En particulier, il existe de nombreux logiciels de planification, la difficulté principale étant de comprendre leurs formalismes puis d'exprimer les problèmes à résoudre dans ces formalismes.

Exécution des adaptations

- Côté entité de contrôle, la fonction d'exécution des adaptations orchestre les appels aux actionneurs, allant de simples appels d'actionneurs sur des paramètres donnés jusqu'à l'utilisation d'un moteur d'exécution de plans.
- Un moteur d'exécution de plans est chargé de veiller :
 - au respect des conditions de déclenchement et à l'atteinte des effets prévus pour chaque appels des actionneurs ;
 - au respect des contraintes temporelles par le déclenchement, la synchronisation et les attentes nécessaires entre les actions élémentaires ;
 - à la détection des échecs (exceptions) en cours ou en fin d'exécution des plans pour déclencher la réparation ou la replanification.
- Moteur d'exécution et actionneurs de l'entité adaptable coopèrent en échangeant les informations et les signaux nécessaires.

Exécution de plans décentralisée

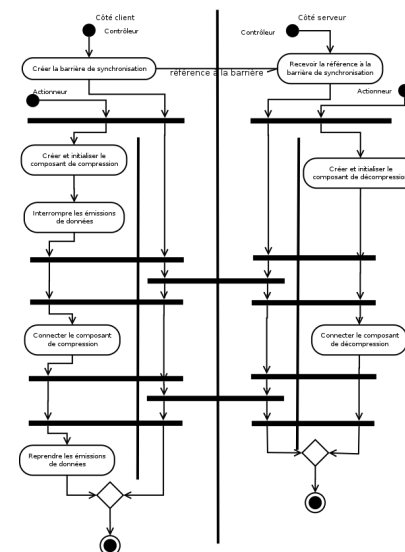
- Le pendant de la planification décentralisée est la *décentralisation de la supervision* entre plusieurs moteurs.
- La supervision décentralisée suppose l'échange de signaux et de données entre moteurs d'exécution (et/ou entre actionneurs), ce qui peut se faire par synchronisation explicite entre ces derniers au prix d'un *couplage fort* entre les deux entités.
- Une autre approche, cherchant à diminuer le couplage, consiste à utiliser des modèles de *coordination* dans cette fonction exécution des entités de contrôle, aspect qui sera abordé au cours 12.
- La *passage à l'échelle* de cette coordination présente cependant des difficultés notoires et va donc nécessiter des techniques novatrices de l'algorithmique répartie ; ce sera un des thèmes abordés au cours 14.

Exemple : retour sur l'exemple de la communication WiFi

- Concernant la synchronisation des adaptations, nous avons présenté au cours 8 une solution fondée sur la coopération directe entre les actionneurs par le partage d'une même barrière de synchronisation fait au moment de leur initialisation.
- Nous présentons maintenant des solutions fondées sur la coopération des moteurs d'exécution des entités de contrôle :
 - en créant dynamiquement la barrière au niveau de l'un ou l'autre des moteurs d'exécution qui s'échangent la référence et la fournissent ensuite à leurs actionneurs respectifs ;
 - en évitant que les deux actionneurs se synchronisent directement l'un avec l'autre en les synchronisant plutôt avec leur moteur d'exécution respectif, moteurs qui se synchronisent entre eux pour assurer indirectement la synchronisation entre les actionneurs.

Choix des entités à synchroniser

Contrôleurs en synchronisation pair-à-pair



Actionneurs en synchronisation pair-à-pair

