
Transfert de Tâches en Apprentissage par Renforcement: Apprentissage Auto-Supervisé de Politiques Généralistes

Louis Bagot
louis.bagot@univ-lyon1.fr

Présentation à MAFTEC + SMAA + ACE
22 Janvier 2026



Objectifs de cette présentation

- Présenter les avancées en Transfert de Tâches en RL
- Tisser des liens avec les autres branches du RL / Deep Learning

Objectifs de cette présentation

- Présenter les avancées en Transfert de Tâches en RL
- Tisser des liens avec les autres branches du RL / Deep Learning
 - 🔗 Exploration et Motivation Intrinsèque
 - 🔗 Apprentissage de compétences et options, RL hiérarchique
 - 🔗 RL Multi-objectif

Objectifs de cette présentation

- Présenter les avancées en Transfert de Tâches en RL
- Tisser des liens avec les autres branches du RL / Deep Learning
 - 🔗 Exploration et Motivation Intrinsèque
 - 🔗 Apprentissage de compétences et options, RL hiérarchique
 - 🔗 RL Multi-objectif
 - 🔗 RL Inversé et Clonage de Politique
 - 🔗 Interprétabilité et Safe RL

Objectifs de cette présentation

- Présenter les avancées en Transfert de Tâches en RL
- Tisser des liens avec les autres branches du RL / Deep Learning
 - 🔗 Exploration et Motivation Intrinsèque
 - 🔗 Apprentissage de compétences et options, RL hiérarchique
 - 🔗 RL Multi-objectif
 - 🔗 RL Inversé et Clonage de Politique
 - 🔗 Interprétabilité et Safe RL
 - 🔗 Apprentissage auto-supervisé

Objectifs de cette présentation

- Présenter les avancées en Transfert de Tâches en RL
- Tisser des liens avec les autres branches du RL / Deep Learning
 - 🔗 Exploration et Motivation Intrinsèque
 - 🔗 Apprentissage de compétences et options, RL hiérarchique
 - 🔗 RL Multi-objectif
 - 🔗 RL Inversé et Clonage de Politique
 - 🔗 Interprétabilité et Safe RL
 - 🔗 Apprentissage auto-supervisé
- faire naître des idées/collaborations?

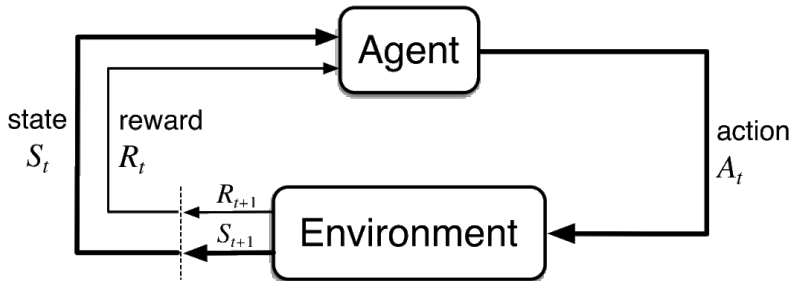
Objectifs de cette présentation

- Présenter les avancées en Transfert de Tâches en RL
- Tisser des liens avec les autres branches du RL / Deep Learning
 - 🔗 Exploration et Motivation Intrinsèque
 - 🔗 Apprentissage de compétences et options, RL hiérarchique
 - 🔗 RL Multi-objectif
 - 🔗 RL Inversé et Clonage de Politique
 - 🔗 Interprétabilité et Safe RL
 - 🔗 Apprentissage auto-supervisé
 - faire naître des idées/collaborations?
- Présenter un peu mon travail (en thèse à IDLab Université d'Anvers)

Déroulé

1. Introduction: Problème de Transfert de Tâches en RL
2. Approche générale en RL Zéro-Shot
3. Familles de Méthodes de Zéro-Shot en RL
4. Méthodes découlant de Successor Features
5. Conclusion

Apprentissage par Renforcement

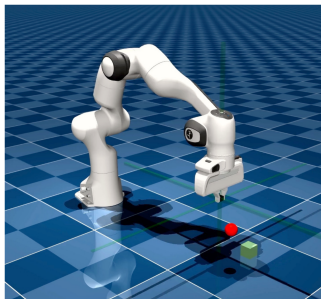


(Sutton et Barto, 2018)

Processus de Décision Markovien (MDP)

$$\text{MDP } \mathcal{M} = \{\mathcal{S}, \mathcal{A}, r, \gamma, p, \mu_0\}$$

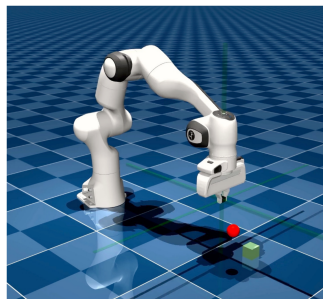
- $\mathcal{S}$: espace d'états
- $\mathcal{A}$: espace d'actions
- $r(s, a)$: fonction de récompense
- $p(s' | s, a)$: fonction de transition
- μ_0 : distribution d'états initiaux



Transfert en RL: **Modification de n'importe quel élément du MDP**

MDP $\mathcal{M} = \{\mathcal{S}, \mathcal{A}, r, \gamma, p, \mu_0\}$

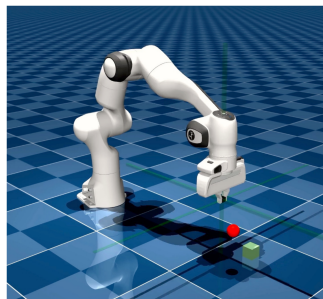
- Δ $\mathcal{S}$: espace d'états
- Δ $\mathcal{A}$: espace d'actions
- Δ $r(s, a)$: fonction de récompense
- Δ $p(s' | s, a)$: fonction de transition
- Δ μ_0 : distribution d'états initiaux



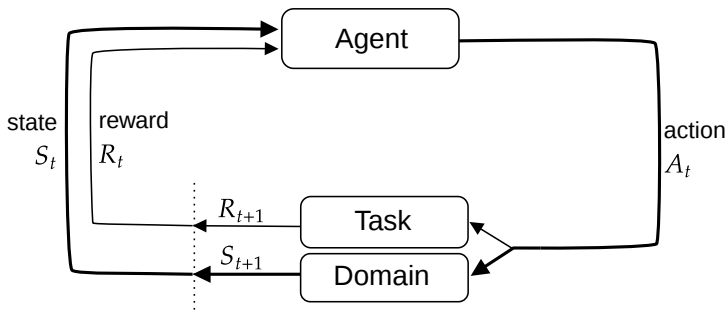
Transfert de Tâche en RL: Modification de la fonction de récompense

$$\text{MDP } \mathcal{M} = \{\mathcal{S}, \mathcal{A}, r, \gamma, p, \mu_0\}$$

- Δ $\mathcal{S}$: espace d'états
- Δ $\mathcal{A}$: espace d'actions
- Δ $r(s, a)$: **fonction de récompense**
- Δ $p(s' | s, a)$: fonction de transition
- Δ μ_0 : distribution d'états initiaux



Tâche et Domaine



Objectif général: maximiser rapidement les récompenses de la nouvelle tâche.

Scénarios de Transfert de Tâches

Scénarios de Transfert de Tâches

- Il y a un lien (ou non) entre les Tâches successives

Scénarios de Transfert de Tâches

- Il y a un lien (ou non) entre les Tâches successives
- La Tâche change doucement (ou subitement)

Scénarios de Transfert de Tâches

- Il y a un lien (ou non) entre les Tâches successives
- La Tâche change doucement (ou subitement)
- L'Agent est au courant (ou non) du changement

Scénarios de Transfert de Tâches

- Il y a un lien (ou non) entre les Tâches successives
- La Tâche change doucement (ou subitement)
- L'Agent est au courant (ou non) du changement
- L'Agent a le droit (ou non) de s'entraîner sur la nouvelle Tâche

Intérêt de cette présentation:
Connaissances sur le Domaine

Intérêt de cette présentation: Connaissances sur le Domaine

- Il y a un lien (**ou non**) entre les Tâches successives
- La Tâche change doucement (**ou subitement**)
- L'Agent est **au courant** (ou non) du changement
- L'Agent a le droit (**ou non**) de s'entraîner sur la nouvelle Tâche

Scénario de *Zéro-shot* RL

Phase de Pré-entraînement:

- L'Agent a accès à un gros dataset d'interactions s, a, s' ; *sans Tâche*
- L'Agent peut s'entraîner dessus pendant longtemps ("offline")

Scénario de Zéro-shot RL

Phase de Pré-entraînement:

- L'Agent a accès à un gros dataset d'interactions s, a, s' ; *sans Tâche*
- L'Agent peut s'entraîner dessus pendant longtemps ("offline")

Phase de Transfert:

- On fournit un dataset de paires état-récompense labélisé (nouvelle Tâche)
- L'Agent doit proposer la meilleure politique possible *sans apprentissage*

Scénario de Zéro-shot RL

Phase de Pré-entraînement:

- L'Agent a accès à un gros dataset d'interactions s, a, s' ; *sans Tâche*
- L'Agent peut s'entraîner dessus pendant longtemps ("offline")

Phase de Transfert:

- On fournit un dataset de paires état-récompense labélisé (nouvelle Tâche)
- L'Agent doit proposer la meilleure politique possible *sans apprentissage*

→ il est tout à fait possible de relaxer des contraintes

Scénario de Zéro-shot RL

Phase de Pré-entraînement:

- L'Agent a accès à un gros dataset d'interactions s, a, s' ; *sans Tâche*
- L'Agent peut s'entraîner dessus pendant longtemps ("offline")

Phase de Transfert:

- On fournit un dataset de paires état-récompense labélisé (nouvelle Tâche)
- L'Agent doit proposer la meilleure politique possible *sans apprentissage*

→ il est tout à fait possible de relaxer des contraintes

 Similarité avec l'apprentissage auto-supervisé!

Déroulé

1. Introduction: Problème de Transfert de Tâches en RL

2. Approche générale en RL Zéro-Shot

2.1. Contexte “historique”: Exploration et Compétences

2.2. Approche générale: conditionnement de politiques

3. Familles de Méthodes de Zéro-Shot en RL

4. Méthodes découlant de Successor Features

5. Conclusion

Overview

1. Introduction: Problème de Transfert de Tâches en RL

2. Approche générale en RL Zéro-Shot

2.1. Contexte “historique”: Exploration et Compétences

2.2. Approche générale: conditionnement de politiques

3. Familles de Méthodes de Zéro-Shot en RL

4. Méthodes découlant de Successor Features

5. Conclusion

Problèmes d'Exploration en RL:

- Exploration vs Exploitation
- Récompenses éparses/tardives: pas de signal, rien à faire
- Exploration intelligente du domaine

Motivation Intrinsèque à la rescousse

Motivation Intrinsèque: générer de fausses récompenses pour motiver l'exploration (Aubret et al., 2023)

- Nouveauté d'état (Bellemare et al., 2016; Burda et al., 2018)
- Curiosité / Surprise du modèle (Burda et al., 2018)
- Progrès d'apprentissage (Oudeyer and Kaplan, 2009)

Motivation Intrinsèque: au delà de l'Exploration

Motivation Intrinsèque: au delà de l'Exploration

Imaginons l'ensemble des fonctions de récompense

Motivation Intrinsèque: au delà de l'Exploration

Imaginons l'ensemble des fonctions de récompense

→ Immensité de comportements/représentations apprenables!

- Atteindre des états ou objectifs
- Eviter des états
- Min/maximiser des features

Motivation Intrinsèque: au delà de l'Exploration

Imaginons l'ensemble des fonctions de récompense

→ Immensité de comportements/représentations apprenables!

- Atteindre des états ou objectifs
- Eviter des états
- Min/maximiser des features

→ Quelles récompenses sont *utiles*?

Motivation Intrinsèque: au delà de l'Exploration

Imaginons l'ensemble des fonctions de récompense

→ Immensité de comportements/représentations apprenables!

- Atteindre des états ou objectifs
- Eviter des états
- Min/maximiser des features

→ Quelles récompenses sont *utiles*?

 Concept similaire aux *Tâches prétexte* en auto-supervisé!

De la Motivation Intrinsèque à l'Apprentissage de Compétences

De la Motivation Intrinsèque à l'Apprentissage de Compétences

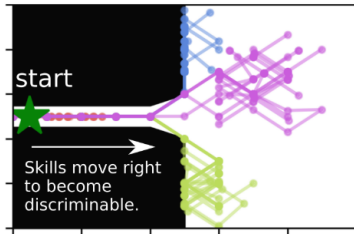
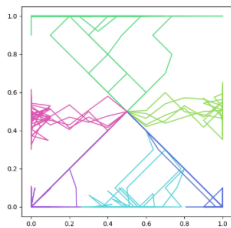
- Est-ce que je peux utiliser plusieurs récompenses intrinsèques?
 - *off-policy*
 - infinité de comportements apprenables depuis s, a, s' !

De la Motivation Intrinsèque à l'Apprentissage de Compétences

- Est-ce que je peux utiliser plusieurs récompenses intrinsèques?
 - *off-policy*
 - infinité de comportements apprenables depuis s, a, s' !
- Quelle *cohésion* entre les récompenses/politiques?
 - éviter la redondance

DIAYN: Apprentissage de compétences via la diversité

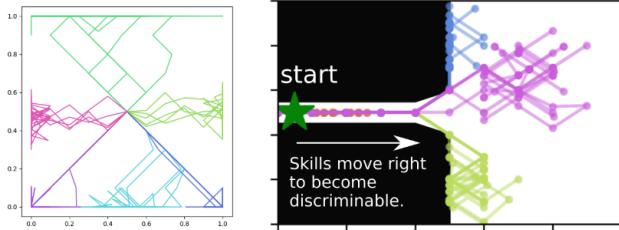
Idée simple “proof of concept”: maximiser la *diversité* (Eysenbach et al. 2018)



→ Notion d'ensemble de *compétences* acquises dans le Domaine!

DIAYN: Apprentissage de compétences via la diversité

Idée simple “proof of concept”: maximiser la *diversité* (Eysenbach et al. 2018)



→ Notion d'ensemble de *compétences* acquises dans le Domaine!

Chaque compétence π_k résout une récompense intrinsèque

→ peut-on les mettre en commun dans $\pi(\cdot, k)$?

Overview

1. Introduction: Problème de Transfert de Tâches en RL

2. Approche générale en RL Zéro-Shot

2.1. Contexte “historique”: Exploration et Compétences

2.2. Approche générale: conditionnement de politiques

3. Familles de Méthodes de Zéro-Shot en RL

4. Méthodes découlant de Successor Features

5. Conclusion

Politiques conditionnées par les Tâches

Zéro-Shot RL: je veux résoudre des Tâches données par une fct récompense r .

Politiques conditionnées par les Tâches

Zéro-Shot RL: je veux résoudre des Tâches données par une fct récompense r .

Idée principale:

1. Trouver une procédure d'*encodage* et *décodage* de r vers $\mathbf{z}_r \in \mathbb{R}^d$

Politiques conditionnées par les Tâches

Zéro-Shot RL: je veux résoudre des Tâches données par une fct récompense r .

Idée principale:

1. Trouver une procédure d'*encodage* et *décodage* de r vers $\mathbf{z}_r \in \mathbb{R}^d$
2. Pré-training: tirer aléatoirement des $\mathbf{z}$ et entraîner $\pi_{\theta}(a \mid s, \mathbf{z})$

Politiques conditionnées par les Tâches

Zéro-Shot RL: je veux résoudre des Tâches données par une fct récompense r .

Idée principale:

1. Trouver une procédure d'*encodage* et *décodage* de r vers $\mathbf{z}_r \in \mathbb{R}^d$
2. Pré-training: tirer aléatoirement des $\mathbf{z}$ et entraîner $\pi_\theta(a \mid s, \mathbf{z})$
 - tirer un $\mathbf{z}$ dans la boule unité (ou exploiter a priori)
 - générer $\hat{r}_z$ à partir de $\mathbf{z}$ avec le décodage
 -  Motivation Intrinsèque!
 - RL: maximiser la récompense (interaction $(s, a, s', \hat{r}_z)$)

Politiques conditionnées par les Tâches

Zéro-Shot RL: je veux résoudre des Tâches données par une fct récompense r .

Idée principale:

1. Trouver une procédure d'*encodage* et *décodage* de r vers $\mathbf{z}_r \in \mathbb{R}^d$
2. Pré-training: tirer aléatoirement des $\mathbf{z}$ et entraîner $\pi_{\theta}(a \mid s, \mathbf{z})$
3. Transfert: compresser r en $\mathbf{z}_r$ et utiliser $\pi_{\theta}(a \mid s, \mathbf{z}_r)$

Liens avec d'autres branches du RL

$$\pi(a \mid s, \mathbf{z})$$

- 🔗 Chaque $\hat{r}_{\mathbf{z}}$ est une récompense intrinsèque
- 🔗 Chaque choix de $\mathbf{z}$ génère une compétence distincte
- 🔗 Hiérarchique: on peut contrôler l'agent avec $\mathbf{z}$!

Est-ce raisonnable?

1. Trouver une procédure d'*encodage* et *décodage* de r vers $\mathbf{z}_r \in \mathbb{R}^d$
2. Pré-training: tirer aléatoirement des $\mathbf{z}$ et entraîner $\pi_{\theta}(a \mid s, \mathbf{z})$
3. Transfert: compresser r en $\mathbf{z}_r$ et utiliser $\pi_{\theta}(a \mid s, \mathbf{z}_r)$

Est-ce raisonnable?

1. Trouver une procédure d'*encodage* et *décodage* de r vers $\mathbf{z}_r \in \mathbb{R}^d$
 2. Pré-training: tirer aléatoirement des $\mathbf{z}$ et entraîner $\pi_{\theta}(a \mid s, \mathbf{z})$
 3. Transfert: compresser r en $\mathbf{z}_r$ et utiliser $\pi_{\theta}(a \mid s, \mathbf{z}_r)$
- la procédure encodage/décodage perd nécessairement de l'information
 - impossible de parfaitement résumer r en général
 - on doit faire des hypothèses sur l'ensemble des fonctions récompenses

Est-ce raisonnable?

1. Trouver une procédure d'encodage et décodage de r vers $\mathbf{z}_r \in \mathbb{R}^d$
 2. Pré-training: tirer aléatoirement des $\mathbf{z}$ et entraîner $\pi_\theta(a \mid s, \mathbf{z})$
 3. Transfert: compresser r en $\mathbf{z}_r$ et utiliser $\pi_\theta(a \mid s, \mathbf{z}_r)$
- la procédure encodage/décodage perd nécessairement de l'information
 - le réseau de neurones $\pi_\theta(a \mid s, \mathbf{z})$ doit emmagasiner une quantité monstrueuse de données et $\mathbf{z}$ (croît avec d)
 - gros réseau, entraîné longtemps, sur beaucoup de données

Est-ce raisonnable?

1. Trouver une procédure d'encodage et décodage de r vers $\mathbf{z}_r \in \mathbb{R}^d$
 2. Pré-training: tirer aléatoirement des $\mathbf{z}$ et entraîner $\pi_\theta(a \mid s, \mathbf{z})$
 3. Transfert: compresser r en $\mathbf{z}_r$ et utiliser $\pi_\theta(a \mid s, \mathbf{z}_r)$
- la procédure encodage/décodage perd nécessairement de l'information
 - le réseau de neurones $\pi_\theta(a \mid s, \mathbf{z})$ doit emmagasiner une quantité monstrueuse de données et $\mathbf{z}$ (croît avec d)
 - gros réseau, entraîné longtemps, sur beaucoup de données
 - généralisation sur tout l'espace de $\mathbf{z}$ compliquée (hors miracles)

Est-ce raisonnable?

1. Trouver une procédure d'encodage et décodage de r vers $\mathbf{z}_r \in \mathbb{R}^d$
 2. Pré-training: tirer aléatoirement des $\mathbf{z}$ et entraîner $\pi_{\theta}(a \mid s, \mathbf{z})$
 3. Transfert: compresser r en $\mathbf{z}_r$ et utiliser $\pi_{\theta}(a \mid s, \mathbf{z}_r)$
- la procédure encodage/décodage perd nécessairement de l'information
 - le réseau de neurones $\pi_{\theta}(a \mid s, \mathbf{z})$ doit emmagasiner une quantité monstrueuse de données et $\mathbf{z}$ (croît avec d)
 - gros réseau, entraîné longtemps, sur beaucoup de données
 - généralisation sur tout l'espace de $\mathbf{z}$ compliquée (hors miracles)
 - redondance: énormément de $\mathbf{z}$ ont les mêmes $\pi(a \mid s, \mathbf{z})$

Degré de liberté de l'approche

1. Trouver une procédure d'*encodage* et *décodage* de r vers $\mathbf{z}_r \in \mathbb{R}^d$
2. Pré-training: tirer aléatoirement des $\mathbf{z}$ et entraîner $\pi(a \mid s, \mathbf{z})$
3. Transfert: compresser r en $\mathbf{z}_r$ et utiliser $\pi(a \mid s, \mathbf{z}_r)$

Degré de liberté de l'approche

1. Trouver une procédure d'*encodage* et *décodage* de r vers $\mathbf{z}_r \in \mathbb{R}^d$
2. Pré-training: tirer aléatoirement des $\mathbf{z}$ et entraîner $\pi(a \mid s, \mathbf{z})$
3. Transfert: compresser r en $\mathbf{z}_r$ et utiliser $\pi(a \mid s, \mathbf{z}_r)$

Degré de liberté de l'approche

1. Trouver une procédure d'*encodage* et *décodage* de r vers $\mathbf{z}_r \in \mathbb{R}^d$
2. Pré-training: tirer aléatoirement des $\mathbf{z}$ et entraîner $\pi(a \mid s, \mathbf{z})$
3. Transfert: compresser r en $\mathbf{z}_r$ et utiliser $\pi(a \mid s, \mathbf{z}_r)$

→ Question critique: **quelle procédure d'encodage/décodage choisir?**

Déroulé

1. Introduction: Problème de Transfert de Tâches en RL

2. Approche générale en RL Zéro-Shot

3. Familles de Méthodes de Zéro-Shot en RL

3.1. Conditionnement par Objectif (Andrychowicz et al. 2017)

3.2. Méthode brute (Frans et al., 2024)

3.3. Successor Features: décomposition linéaire de r (Barreto et al., 2017)

4. Méthodes découlant de Successor Features

5. Conclusion

Overview

1. Introduction: Problème de Transfert de Tâches en RL

2. Approche générale en RL Zéro-Shot

3. Familles de Méthodes de Zéro-Shot en RL

3.1. Conditionnement par Objectif (Andrychowicz et al. 2017)

3.2. Méthode brute (Frans et al., 2024)

3.3. Successor Features: décomposition linéaire de r (Barreto et al., 2017)

4. Méthodes découlant de Successor Features

5. Conclusion

Cas spécial du Conditionnement par Objectif (*goal*)

Si on choisit $\mathbf{z} = g \in \mathcal{S}$ (ou espace latent d'états), $\pi(a \mid s, g)$

- π cherche simplement à *atteindre* g

Cas spécial du Conditionnement par Objectif (*goal*)

Si on choisit $\mathbf{z} = g \in \mathcal{S}$ (ou espace latent d'états), $\pi(a \mid s, g)$

- π cherche simplement à *atteindre* g
- Récompense sparse mais pas un problème: (Andrychowicz et al. 2017)
 - pour chaque transition (s, a, s') , on peut prétendre que $g = s'$

Cas spécial du Conditionnement par Objectif (*goal*)

Si on choisit $\mathbf{z} = g \in \mathcal{S}$ (ou espace latent d'états), $\pi(a \mid s, g)$

- π cherche simplement à *atteindre* g
- Récompense sparse mais pas un problème: (Andrychowicz et al. 2017)
 - pour chaque transition (s, a, s') , on peut prétendre que $g = s'$
 - toute transition peut être une réussite!
 - même idée que “toute transition peut être vue du pdv de toute récompense”

Cas spécial du Conditionnement par Objectif (*goal*)

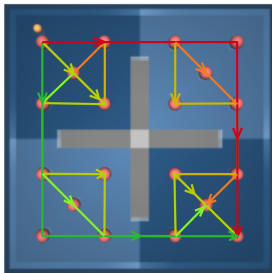
Si on choisit $\mathbf{z} = g \in \mathcal{S}$ (ou espace latent d'états), $\pi(a \mid s, g)$

- π cherche simplement à *atteindre* g
- Pas vraiment du zero-shot
 - “encodage”: $g = \arg \max_s r(s)$
 - “décodage”: $\hat{r}(s) = \mathbb{1}_{s=g}$ (ou pareil -1)

Cas spécial du Conditionnement par Objectif (*goal*)

Si on choisit $\mathbf{z} = g \in \mathcal{S}$ (ou espace latent d'états), $\pi(a \mid s, g)$

- π cherche simplement à *atteindre* g
- Pas vraiment du zero-shot
 - “encodage”: $g = \arg \max_s r(s)$
 - “décodage”: $\hat{r}(s) = \mathbb{1}_{s=g}$ (ou pareil -1)
- plein de trajectoires peuvent mener au même état objectif



Cas spécial du Conditionnement par Objectif (*goal*)

Si on choisit $\mathbf{z} = g \in \mathcal{S}$ (ou espace latent d'états), $\pi(a \mid s, g)$

- π cherche simplement à *atteindre* g
- Pas vraiment du zero-shot
 - "encodage": $g = \arg \max_s r(s)$
 - "décodage": $\hat{r}(s) = \mathbb{1}_{s=g}$ (ou pareil -1)
- formulation récente en contrastif (Eysenbach et al., 2022, Wang et al., 2025)
 - 🔗 lien avec l'auto-supervisé!

Overview

1. Introduction: Problème de Transfert de Tâches en RL

2. Approche générale en RL Zéro-Shot

3. Familles de Méthodes de Zéro-Shot en RL

3.1. Conditionnement par Objectif (Andrychowicz et al. 2017)

3.2. Méthode brute (Frans et al., 2024)

3.3. Successor Features: décomposition linéaire de r (Barreto et al., 2017)

4. Méthodes découlant de Successor Features

5. Conclusion

Utiliser un réseau pour faire l'encodage/décodage

Version la plus brute: *Functional Reward Encodings* (Frans et al., 2024)

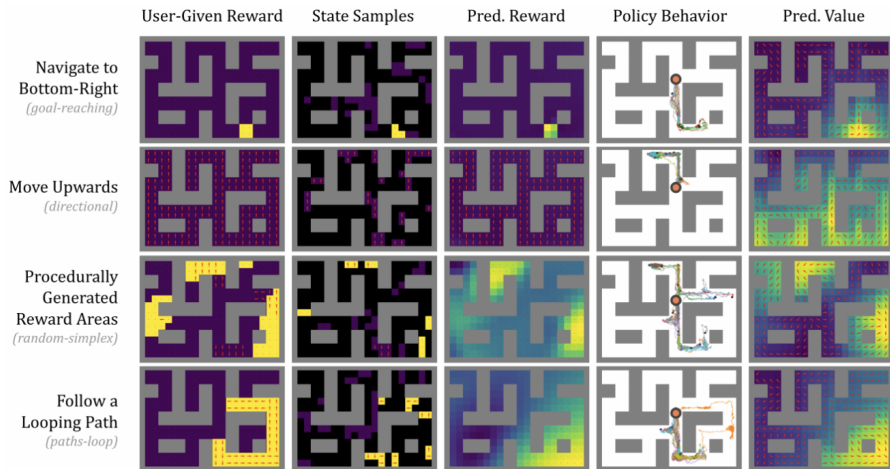
- réseau encodeur $\xi_{\theta}(\{s_i, r(s_i)\}_i) = \mathbf{z}_r$
- réseau décodeur $\xi_{\mathbf{w}}^{-1}(\{\mathbf{z}_r, s_i\}_i) = \{\hat{r}(s_i)\}_i$

Utiliser un réseau pour faire l'encodage/décodage

Version la plus brute: *Functional Reward Encodings* (Frans et al., 2024)

- réseau encodeur $\xi_{\theta}(\{s_i, r(s_i)\}_i) = \mathbf{z}_r$
- réseau décodeur $\xi_w^{-1}(\{\mathbf{z}_r, s_i\}_i) = \{\hat{r}(s_i)\}_i$
- setup le plus flexible
- boîte noire absolue
- pas évident de tirer aléatoirement des $\mathbf{z}$ pour l'entraînement

Functional Reward Encodings: visualisation



Overview

1. Introduction: Problème de Transfert de Tâches en RL

2. Approche générale en RL Zéro-Shot

3. Familles de Méthodes de Zéro-Shot en RL

3.1. Conditionnement par Objectif (Andrychowicz et al. 2017)

3.2. Méthode brute (Frans et al., 2024)

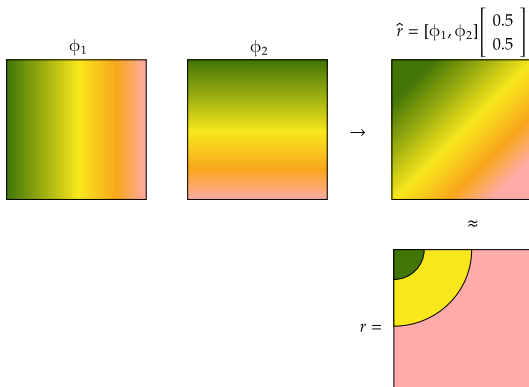
3.3. Successor Features: décomposition linéaire de r (Barreto et al., 2017)

4. Méthodes découlant de Successor Features

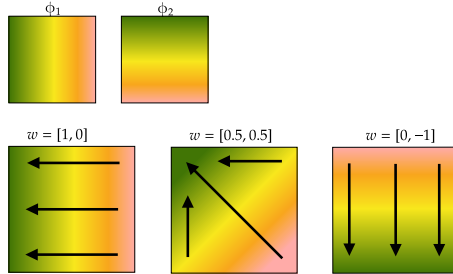
5. Conclusion

Successor Features: Intuition (Barreto et al., 2017, 2018)

Approximation linéaire de la récompense avec *features* $\{\phi_i\}_i$: $r \approx \phi \cdot \mathbf{z}_r$

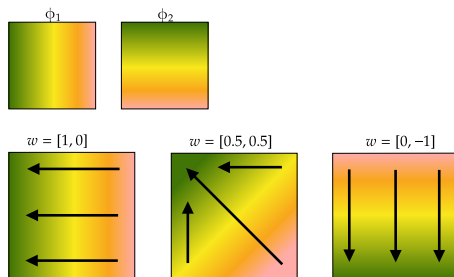


Conditionnement sur le vecteur d'approx. linéaire



Observation: pour une base ϕ donnée, **z encode la récompense.**

Conditionnement sur le vecteur d'approx. linéaire



Observation: pour une base ϕ donnée, **$\mathbf{z}$ encode la récompense.**
→ Entraîner $\pi(a \mid s, \mathbf{z})$.

Successor Features comme procédure encodage/décodage

Pour une base des récompenses $\phi : \mathcal{S} \rightarrow \mathbb{R}^d$,

- Encodage: moindres carrés $\xi(\{s_i, r(s_i)\}_i) = \mathbf{z}_r = \arg \min_{\mathbf{z}} \|r - \phi \cdot \mathbf{z}\|_2$
- Décodage: approx. lin. $\xi^{-1}(s, \mathbf{z}_r) = \phi(s) \cdot \mathbf{z}_r$

Successor Features comme procédure encodage/décodage

Pour une base des récompenses $\phi : \mathcal{S} \rightarrow \mathbb{R}^d$,

- Encodage: moindres carrés $\xi(\{s_i, r(s_i)\}_i) = \mathbf{z}_r = \arg \min_{\mathbf{z}} \|\mathbf{r} - \phi \cdot \mathbf{z}\|_2$
- Décodage: approx. lin. $\xi^{-1}(s, \mathbf{z}_r) = \phi(s) \cdot \mathbf{z}_r$
- Chaque z_i correspond à la **contribution du feature** ϕ_i
- $\pi(a \mid s, \mathbf{z})$ apprend à maximiser les comb. lin. de features

Parallèles avec le RL Multi-Objectif

En RL Multi Objectif:

- Ensemble de récompenses $\mathbf{r} = [r_0, r_1, \dots, r_d]$
- Version linéaire: maximiser une comb. lin. d'objectifs $\mathbf{r} \cdot \mathbf{z}$



Parallèles avec le RL Multi-Objectif

En RL Multi Objectif:

- Ensemble de récompenses $\mathbf{r} = [r_0, r_1, \dots, r_d]$
 - Version linéaire: maximiser une comb. lin. d'objectifs $\mathbf{r} \cdot \mathbf{z}$
- Même setup, mais en SF on choisit l'ensemble de récompenses ϕ (eq. $\mathbf{r}$) (récompenses intrinsèques)

Parallèles avec le RL Multi-Objectif

En RL Multi Objectif:

- Ensemble de récompenses $\mathbf{r} = [r_0, r_1, \dots, r_d]$
 - Version linéaire: maximiser une comb. lin. d'objectifs $\mathbf{r} \cdot \mathbf{z}$
- Même setup, mais en SF on choisit l'ensemble de récompenses ϕ (eq. $\mathbf{r}$) (récompenses intrinsèques)

Même notion de fonction de valeur vectorielle:

- Multi-Objectif: $\mathbf{V}^\pi = \mathbb{E}_\pi[\sum_t \mathbf{r}(S_t)]$
- Définition de *Successor Features*: $\boldsymbol{\psi}^\pi = \mathbb{E}_\pi[\sum_t \phi(S_t)]$
- Et pour une comb. lin. $\mathbf{z}$, $v_z^\pi = \mathbf{V}^\pi \cdot \mathbf{z}$ (idem $\boldsymbol{\psi}$)

Successor Features comme procédure encodage/décodage

Pour une base des récompenses $\phi : \mathcal{S} \rightarrow \mathbb{R}^d$,

- Chaque z_i correspond à la **contribution du feature** ϕ_i
- $\pi(a \mid s, \mathbf{z})$ apprend à maximiser les comb. lin. de features

Successor Features comme procédure encodage/décodage

Pour une base des récompenses $\phi : \mathcal{S} \rightarrow \mathbb{R}^d$,

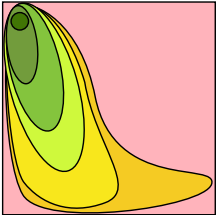
- Chaque z_i correspond à la **contribution du feature** ϕ_i
- $\pi(a \mid s, \mathbf{z})$ apprend à maximiser les comb. lin. de features
- Beaucoup plus interprétable
- Beaucoup moins flexible

Successor Features comme procédure encodage/décodage

Pour une base des récompenses $\phi : \mathcal{S} \rightarrow \mathbb{R}^d$,

- Chaque z_i correspond à la **contribution du feature** ϕ_i
- $\pi(a \mid s, \mathbf{z})$ apprend à maximiser les comb. lin. de features
- Beaucoup plus interprétable
- Beaucoup moins flexible
- ? **Quel choix de features/base linéaire ϕ ?**

Quel choix de ϕ pour approximer r linéairement?

$$r \approx \phi \cdot w$$

$$\approx \begin{array}{c} \boxed{?} \times w_1 \\ + \\ \boxed{?} \times w_2 \\ + \\ \boxed{?} \times w_3 \\ + \\ \boxed{?} \times w_4 \\ + \\ \boxed{?} \times w_5 \\ \vdots \end{array} = \boxed{?}$$

Déroulé

1. Introduction: Problème de Transfert de Tâches en RL

2. Approche générale en RL Zéro-Shot

3. Familles de Méthodes de Zéro-Shot en RL

4. Méthodes découlant de Successor Features

4.1. Successor Clusters (Bagot et al., 2025)

4.2. Hilbert Foundation Policies (HILP, Park et al. 2024)

4.3. Fonctions propres du Laplacien (Machado et al., 2023)

4.4. Applications intéressantes de ces méthodes

5. Conclusion

Rappel du problème

Rappel du problème

- Objectif: transfert zéro-shot
 - extraire des connaissances du domaine pour résoudre toute Tâche

Rappel du problème

- Objectif: transfert zéro-shot
 - extraire des connaissances du domaine pour résoudre toute Tâche
- ↓
- Politiques conditionnées par la Tâche $\pi(a \mid s, \mathbf{z})$
 - besoin d'une procédure d'encodage/décodage de r vers $\mathbf{z}_r$

Rappel du problème

- Objectif: transfert zéro-shot

→ extraire des connaissances du domaine pour résoudre toute Tâche



- Politiques conditionnées par la Tâche $\pi(a \mid s, \mathbf{z})$

→ besoin d'une procédure d'encodage/décodage de r vers $\mathbf{z}_r$



- Successor Features: approximation linéaire de récompense $r \approx \boldsymbol{\phi} \cdot \mathbf{z}_r$

→ $\mathbf{z}_r$ encode l'importance de chaque feature ϕ_i dans r

Rappel du problème

- Objectif: transfert zéro-shot

→ extraire des connaissances du domaine pour résoudre toute Tâche



- Politiques conditionnées par la Tâche $\pi(a \mid s, \mathbf{z})$

→ besoin d'une procédure d'encodage/décodage de r vers $\mathbf{z}_r$



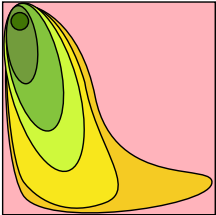
- Successor Features: approximation linéaire de récompense $r \approx \phi \cdot \mathbf{z}_r$

→ $\mathbf{z}_r$ encode l'importance de chaque feature ϕ_i dans r



- Quel choix de features/base linéaire ϕ ?

Quel choix de ϕ pour approximer r linéairement?

$$r \approx \phi \cdot w$$

$$\approx \begin{array}{c} \boxed{?} \times w_1 \\ + \\ \boxed{?} \times w_2 \\ + \\ \boxed{?} \times w_3 \\ + \\ \boxed{?} \times w_4 \\ + \\ \boxed{?} \times w_5 \\ \vdots \end{array} = \boxed{?}$$

Overview

1. Introduction: Problème de Transfert de Tâches en RL

2. Approche générale en RL Zéro-Shot

3. Familles de Méthodes de Zéro-Shot en RL

4. Méthodes découlant de Successor Features

4.1. Successor Clusters (Bagot et al., 2025)

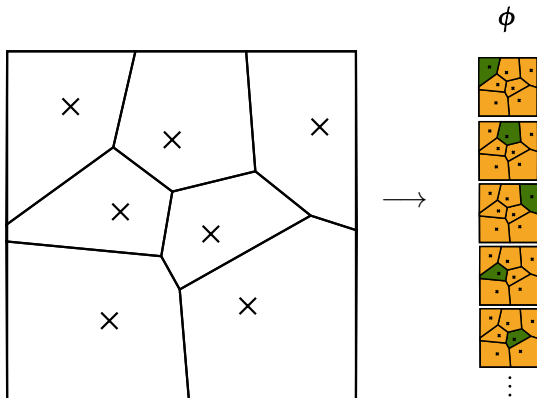
4.2. Hilbert Foundation Policies (HILP, Park et al. 2024)

4.3. Fonctions propres du Laplacien (Machado et al., 2023)

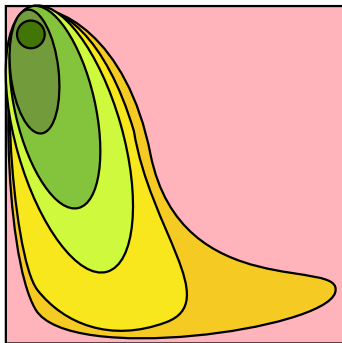
4.4. Applications intéressantes de ces méthodes

5. Conclusion

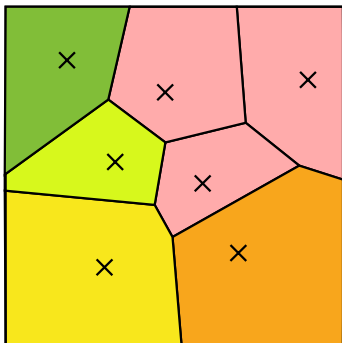
Occupation de cluster comme features



Approximation de récompense via clusters



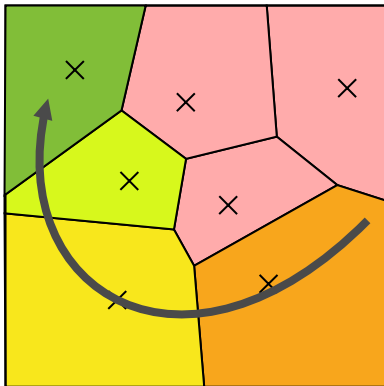
true reward r



reward approximation

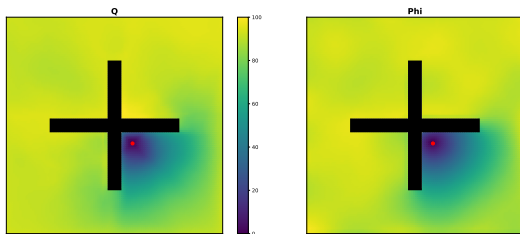
Politique conditionnée par les valeurs des clusters

$\pi(a | s, \mathbf{z})$ conditionnée sur $\mathbf{z}$ avec z_i = récompense moyenne du cluster i

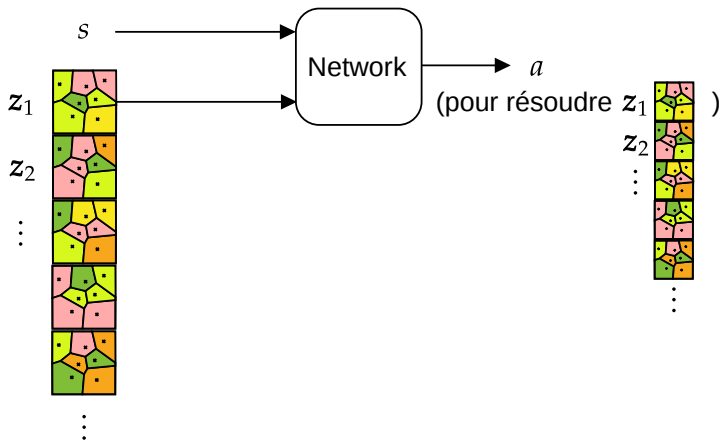


Notion de distance pour les clusters: Distance temporelle minimale

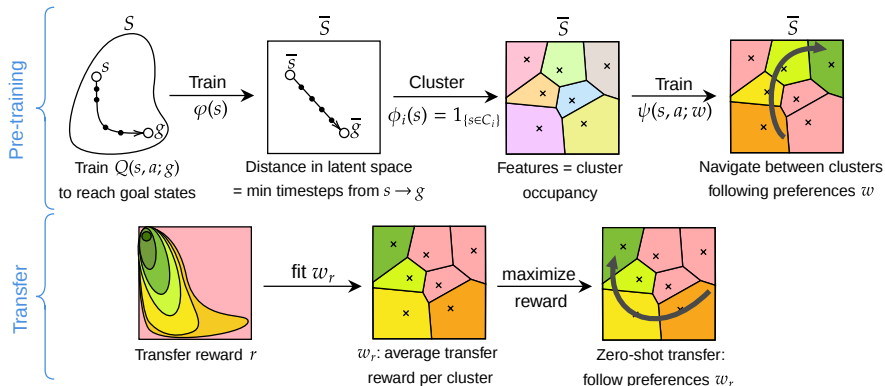
1. Entraîner un agent conditionné par l'objectif $\pi(a \mid s, g)$
2. Récompense prétexte = $-\mathbb{1}_{s \neq g}$
→ fonction de valeur = distance temporelle minimale
3. Apprendre un espace latent où dist. eudlidienne = dist. temporelle



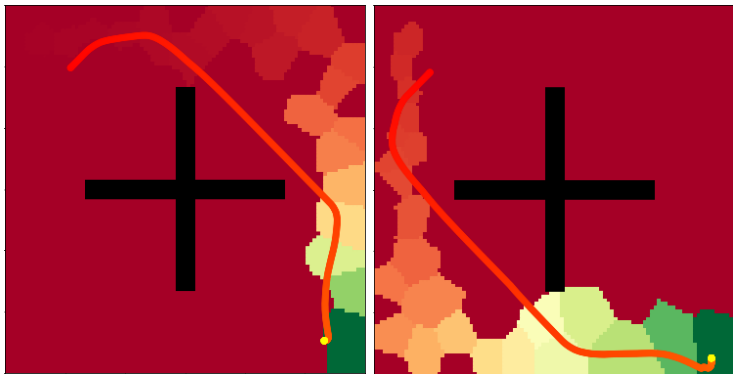
Entraînement du réseau: génération de z prétexte



Méthode complète: RL Zéro-shot avec features de clusters



Contrôle de l'agent via z



background: ce que l'agent prend en entrée
(vecteur 100 dimensions pour 100 clusters)

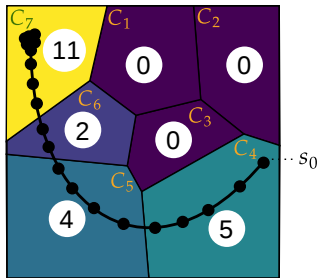
Successor Clusters: prédiction de visites futures de clusters

On a établi que $\psi^\pi(s_0) = \mathbb{E}_\pi[\sum_t \phi(s_t) \mid s_0]$.

$$\psi^\pi(s_0) = \phi(s_0) + \phi(s_1) + \phi(s_2) + \phi(s_3) + \dots$$

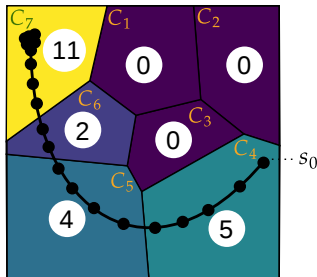


Successor Clusters sur tout une trajectoire



$$\psi^\pi(s_0) = \begin{bmatrix} 0 \\ 0 \\ 0 \\ 5 \\ 4 \\ 2 \\ 11 \end{bmatrix}$$

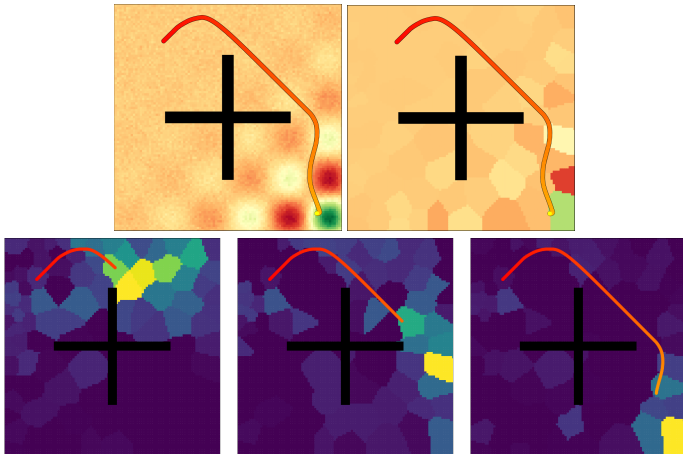
Successor Clusters sur tout une trajectoire



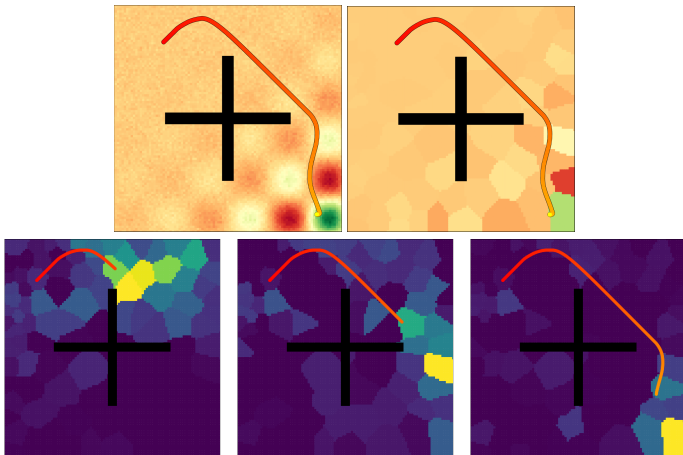
$$\psi^\pi(s_0) = \begin{bmatrix} 0 \\ 0 \\ 0 \\ 5 \\ 4 \\ 2 \\ 11 \end{bmatrix}$$

En général, $\psi(s_0, a_0, \mathbf{z}) = \mathbb{E}_{\pi(\cdot|\mathbf{z})}[\sum_t \phi(s_t) \mid s_0, a_0]$
(USFA, Borsa et al., 2018)

Successor Clusters: Exemple Zero-shot



Successor Clusters: Exemple Zero-shot



 Interprétabilité: l'agent prédit sa trajectoire future.

Overview

1. Introduction: Problème de Transfert de Tâches en RL

2. Approche générale en RL Zéro-Shot

3. Familles de Méthodes de Zéro-Shot en RL

4. Méthodes découlant de Successor Features

4.1. Successor Clusters (Bagot et al., 2025)

4.2. Hilbert Foundation Policies (HILP, Park et al. 2024)

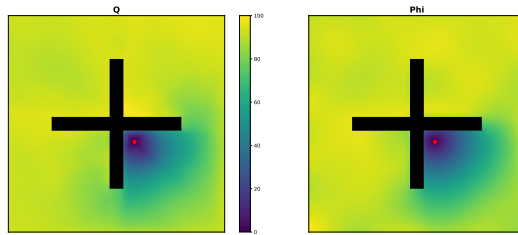
4.3. Fonctions propres du Laplacien (Machado et al., 2023)

4.4. Applications intéressantes de ces méthodes

5. Conclusion

HILP: Autre utilisation de l'espace

"dist. euclidienne = dist. temporelle"

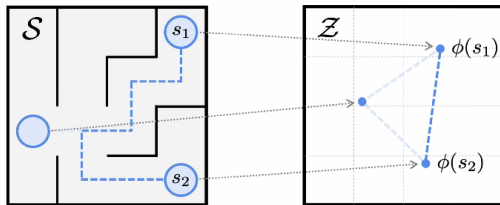


HILP: Autre utilisation de l'espace

"dist. euclidienne = dist. temporelle"

- Quasiment le même espace latent que Successor Clusters

Hilbert representation $\phi(s)$



Distance-preserving mapping $d^*(s_1, s_2) = \|\phi(s_1) - \phi(s_2)\|$

HILP: Autre utilisation de l'espace

"dist. euclidienne = dist. temporelle"

- Idée clé: **tirer z directement de l'espace latent!**

HILP: Autre utilisation de l'espace

"dist. euclidienne = dist. temporelle"

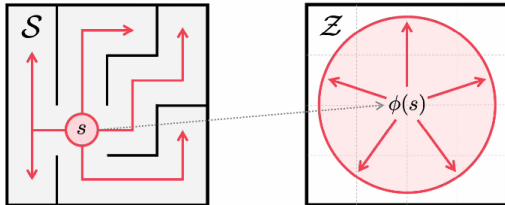
- Idée clé: **tirer $\mathbf{z}$ directement de l'espace latent!**
- Tout $\mathbf{z}$ (normalisé) correspond à une *direction* dans l'espace latent

HILP: Autre utilisation de l'espace "dist. euclidienne = dist. temporelle"

- Idée clé: **tirer $\mathbf{z}$ directement de l'espace latent!**
 - Tout $\mathbf{z}$ (normalisé) correspond à une *direction* dans l'espace latent
- récompense pour suivre cette direction $\mathbf{z}$:

$$\text{décodeur } r_z(s, a, s') = (\varphi(s') - \varphi(s)) \cdot \mathbf{z}$$

Hilbert foundation policy $\pi(a | s, \mathbf{z})$



Directional state-spanning policies $r = \langle \phi(s') - \phi(s), \mathbf{z} \rangle$

HILP: Autre utilisation de l'espace "dist. euclidienne = dist. temporelle"

- Idée clé: **tirer $\mathbf{z}$ directement de l'espace latent!**
 - Tout $\mathbf{z}$ (normalisé) correspond à une *direction* dans l'espace latent
- récompense pour suivre cette direction $\mathbf{z}$:
- $$\text{décodeur } r_z(s, a, s') = (\varphi(s') - \varphi(s)) \cdot \mathbf{z}$$
- Entraîner la politique $\pi(a \mid s, \mathbf{z})$ avec $\mathbf{z} \in \mathcal{Z} \subset \mathbb{R}^d$, l'espace latent
 - la base orthonormale de $\mathcal{Z}$ définit nos récompenses intrinsèques

HILP: Autre utilisation de l'espace "dist. euclidienne = dist. temporelle"

- Idée clé: **tirer $\mathbf{z}$ directement de l'espace latent!**
 - Tout $\mathbf{z}$ (normalisé) correspond à une *direction* dans l'espace latent
- récompense pour suivre cette direction $\mathbf{z}$:
- $$\text{décodeur } r_z(s, a, s') = (\varphi(s') - \varphi(s)) \cdot \mathbf{z}$$
- Entraîner la politique $\pi(a \mid s, \mathbf{z})$ avec $\mathbf{z} \in \mathcal{Z} \subset \mathbb{R}^d$, l'espace latent
 - la base orthonormale de $\mathcal{Z}$ définit nos récompenses intrinsèques
 - On peut faire du conditionnement d'objectif: $\mathbf{z}_g(s) = \varphi(g) - \varphi(s)$
 - vecteur de Tâche $\mathbf{z}$ *dynamique*!

Overview

1. Introduction: Problème de Transfert de Tâches en RL

2. Approche générale en RL Zéro-Shot

3. Familles de Méthodes de Zéro-Shot en RL

4. Méthodes découlant de Successor Features

4.1. Successor Clusters (Bagot et al., 2025)

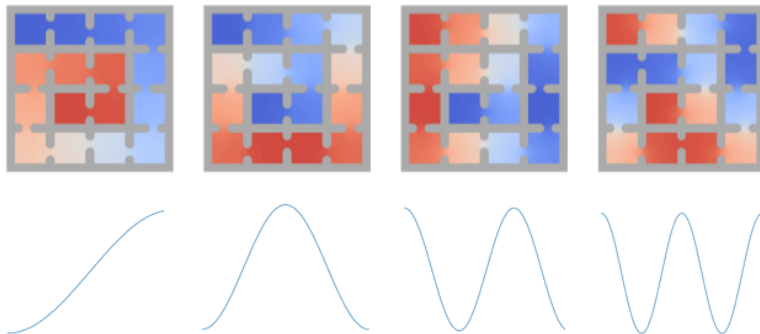
4.2. Hilbert Foundation Policies (HILP, Park et al. 2024)

4.3. Fonctions propres du Laplacien (Machado et al., 2023)

4.4. Applications intéressantes de ces méthodes

5. Conclusion

Fonctions propres du Laplacien (K. Wang et al., 2021)



Equivalent MDP d'une base de Fourier → très adaptées à une combinaison linéaire!

Exploration et Zéro-Shot RL avec le Laplacien

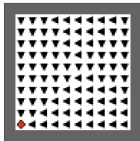
- Pour f_i chaque fonction propre on utilise $\phi = [f_1, f_2 \dots f_d]$
 - les f_i servent de motivation intrinsèque (Machado et al., 2023)

Exploration et Zéro-Shot RL avec le Laplacien

- Pour f_i chaque fonction propre on utilise $\phi = [f_1, f_2 \dots f_d]$
 - les f_i servent de motivation intrinsèque (Machado et al., 2023)
- z_i correspond à la contribution de chaque f_i
- On entraîne $\pi(a \mid s, \mathbf{z})$.

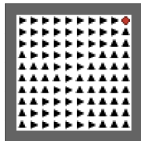
Visualisation de l'agent "Laplacien"

1st eigenoption



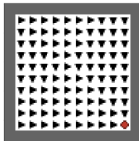
[1,0,0,1,0,0,0,0,0,0]

2nd eigenoption



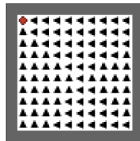
[1,0,1,0,0,0,0,0,0,0]

3rd eigenoption



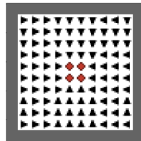
[0,1,0,1,0,0,0,0,0,0]

4th eigenoption

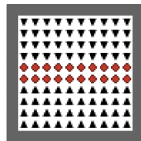
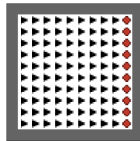
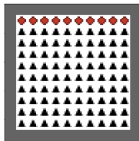
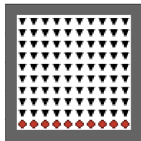
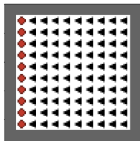


[0,1,1,0,0,0,0,0,0,0]

10th eigenoption



[0,0,0,0,0,0,1,0,0,1]



Autre méthode intéressante: Forward-Backward Representations

(Touati et al., 2021, 2023)

pas le temps désolé!

Overview

1. Introduction: Problème de Transfert de Tâches en RL

2. Approche générale en RL Zéro-Shot

3. Familles de Méthodes de Zéro-Shot en RL

4. Méthodes découlant de Successor Features

4.1. Successor Clusters (Bagot et al., 2025)

4.2. Hilbert Foundation Policies (HILP, Park et al. 2024)

4.3. Fonctions propres du Laplacien (Machado et al., 2023)

4.4. Applications intéressantes de ces méthodes

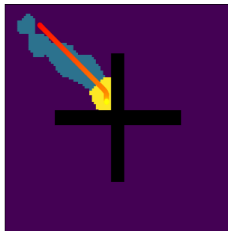
5. Conclusion

Inverse RL / Clonage de Politique

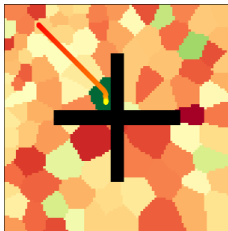
- Pour une trajectoire désirée $S_0, S_1, S_2, \dots$
- Somme de features désirée $I = \sum_t \phi(S_t)$
- Optimiser $\mathbf{z}$ pour minimiser $\|I - \psi(S_0, \mathbf{z})\|$

Inverse RL / Clonage de Politique

- Pour une trajectoire désirée $S_0, S_1, S_2, \dots$
- Somme de features désirée $\Gamma = \sum_t \phi(S_t)$
- Optimiser $\mathbf{z}$ pour minimiser $\|\Gamma - \psi(S_0, \mathbf{z})\|$



trajectoire désirée Γ ;



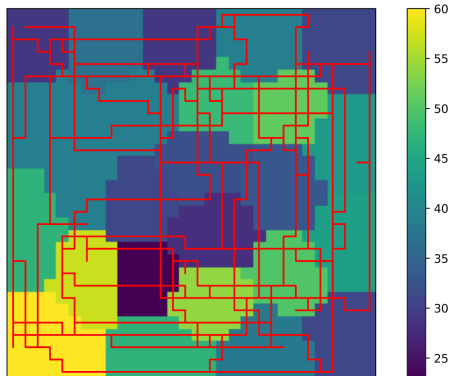
$\mathbf{z}_\Gamma$ optimisé;



politique et visites pour $\mathbf{z}_\Gamma$

📌 Générer des politiques d'Exploration

Clusters: compter et punir chaque visite $\rightarrow \mathbf{z} = 1/[N(C_1), N(C_2), \dots, N(C_d)]$



→ Généralisation à venir...

Safe RL et Interprétabilité

Entraînement: on peut contrôler les features ϕ_i et $\mathbf{z}$

- Ajouter à la main des features/régions d'états dangereux
- Entraîner seulement pour les $\mathbf{z}$ qui nous intéressent
 - pas besoin de maximiser les features dangereux!

Safe RL et Interprétabilité

Entraînement: on peut contrôler les features ϕ_i et $\mathbf{z}$

- Ajouter à la main des features/régions d'états dangereux
- Entraîner seulement pour les $\mathbf{z}$ qui nous intéressent
 - pas besoin de maximiser les features dangereux!

Transfert/exécution: on peut contrôler les visites prévues ψ_i

- Stopper l'agent en cas de prédiction trop élevée
- Prédiction encore très imparfaite

Déroulé

1. Introduction: Problème de Transfert de Tâches en RL
2. Approche générale en RL Zéro-Shot
3. Familles de Méthodes de Zéro-Shot en RL
4. Méthodes découlant de Successor Features
- 5. Conclusion**

Résumé: méthodes abordées

- On veut extraire des politique et représentations du domaine pour le transfert
- Politique conditionnée: $\pi(a \mid s, \mathbf{z})$
- Dépend d'une méthode d'encodage des tâches $r \rightarrow \mathbf{z}$
 - Conditionnement d'objectifs: $\mathbf{z}$ est un état
 - Encodage brut à base de réseau: flexible mais obscur
 - Successor Features: décomposition linéaire

Résumé: liens tissés

- 🔗 Exploration & Motivation: *ensemble* de récompenses intrinsèques
- 🔗 Compétences & Hiérarchique: politiques fortes et contrôlables avec $\mathbf{z}$
 - problème: pas pensées pour être composées
- 🔗 Multi-Objectif: même scénario, mais on choisit les récompenses
- 🔗 RL Inversé & Clonage: on peut extraire le $\mathbf{z}$ qui suit une trajectoire
- 🔗 RL Safe & Interprétabilité: l'agent prédit ses trajectoires
- 🔗 Apprentissage auto-supervisé: on réfléchit à des tâches prétexte

Limites et futur

Plein de facteurs limitants:

- Objectif de $\pi(a \mid s, \mathbf{z})$ très compliqué
 - surtout vu les redondances dans $\mathbf{z}$
- Pas de compositionnalité
 - π essaie de tout résoudre seule!
- Setup de zero-shot trop simplifié
 - pas d'exploration
 - gros dataset & temps d'entraînement
 - transfert idéalisé

Merci pour votre attention!

Appendix: éléments complémentaires

Déroulé

6. Successor Features & Clusters

- 6.1. Deeper dive into the mathematics
- 6.2. Visualizations and Experiments

7. Future work: concrete propositions

Overview

6. Successor Features & Clusters

6.1. Deeper dive into the mathematics

6.2. Visualizations and Experiments

7. Future work: concrete propositions

Successor Features

Approximate reward

$$r \approx \phi \cdot \mathbf{w}_r$$

Successor Features are q -values for the features ϕ_i :

$$\psi_{\pi}(s, a) = \mathbb{E} \left[\sum_{k=0}^{\infty} \gamma^k \phi(S_{t+k+1}, A_{t+k+1}) \mid S_t = s, A_t = a \right].$$

Successor Features linearly break down the expected return:

$$\begin{aligned} q_{\pi}(s, a) &= \mathbb{E} \left[\sum_{k=0}^{\infty} \gamma^k R_{t+k+1} \mid S_t = s, A_t = a \right] \\ &\approx \psi_{\pi}(s, a) \cdot \mathbf{w}_r. \end{aligned}$$

Generating Optimal Policies with USFAs

Compute USFAs with Q learning:

$$\psi(s, a; \mathbf{w}_r) \leftarrow \phi(s, a) + \gamma \max \psi(s', a^*; \mathbf{w}_r).$$

Obtain optimal policy with

$$\pi(s, \mathbf{w}_r) = \arg \max_a \psi(s, a; \mathbf{w}_r) \cdot \mathbf{w}_r.$$

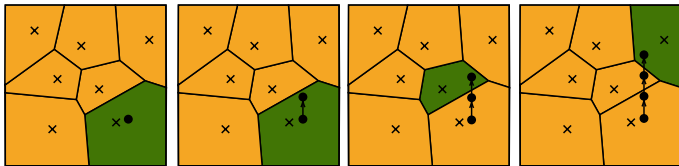
Successor Clusters predict future cluster visits

- Linear reward approximation: $r \approx \phi \cdot \mathbf{w}_r$.
- Value function breaks down linearly: $q_\pi \approx \psi_\pi \cdot \mathbf{w}_r$

→ with $\psi_{\pi,i}$ the value function for ϕ_i .

(Dayan, 1993; Barreto et al., 2017, 2018)

$$\psi^\pi(s_0) = \phi(s_0) + \phi(s_1) + \phi(s_2) + \phi(s_3) + \dots$$



Overview

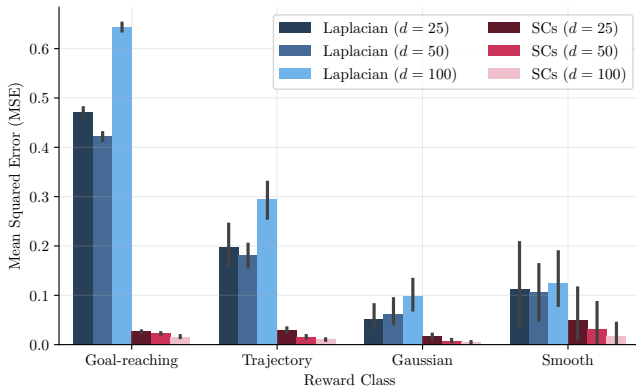
6. Successor Features & Clusters

6.1. Deeper dive into the mathematics

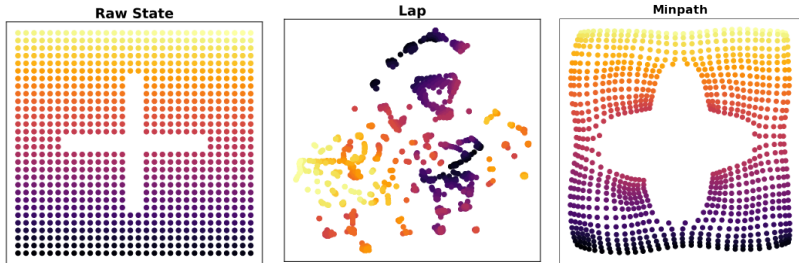
6.2. Visualizations and Experiments

7. Future work: concrete propositions

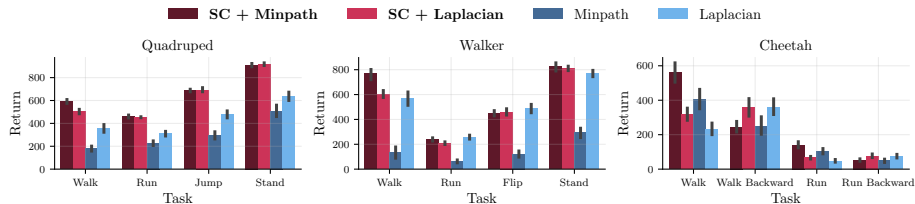
Reward Approximation



Latent space comparison

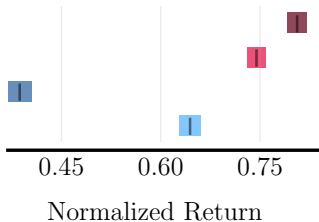


Zero-shot RL results



Overall Performance

SC + Minpath
SC + Laplacian
Minpath
Laplacian



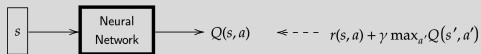
Déroulé

6. Successor Features & Clusters

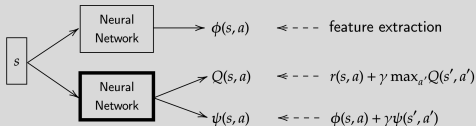
7. Future work: concrete propositions

Between single-task and task-conditioned policies

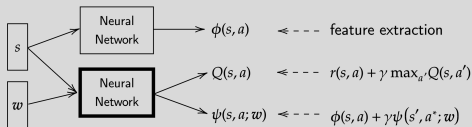
Traditional RL Agent (e.g. DQN)



Fine-Tuner Agent



Task Transfer Agent



Proof of concept “fine-tuner” agent

