

Recommandation de tests multi-objectifs pour l'apprentissage adaptatif

Nassim Bouarour¹ Idir Benouaret² Sihem Amer-Yahia¹

(1) LIG, 700 Av. Centrale, 38401 Saint-Martin-d'Hères, France

(2) LRE, 14-16 Rue Voltaire, 94270 Le Kremlin-Bicêtre, France

nassim.bouarour@univ-grenoble-alpes.fr, idir.benouaret@epita.fr,
sihem.amer-yahia@univ-grenoble-alpes.fr

RÉSUMÉ

L'amélioration des compétences (*upskilling*) est un segment en forte croissance en éducation. Pourtant, peu de travaux algorithmiques se concentrent sur l'élaboration de stratégies dédiées pour atteindre une maîtrise avancée des compétences. Dans cet article, nous formalisons AdUp, un problème d'amélioration itérative des compétences combinant l'apprentissage par maîtrise et la théorie de la Zone de Développement Proximal. Nous étendons nos travaux précédents et concevons deux solutions pour AdUp : MOO et MAB. MOO est une approche d'optimisation multi-objectifs qui utilise une méthode de *Hill Climbing* pour adapter la difficulté des tests recommandés selon 3 objectifs : la performance prédite de l'apprenant, son aptitude, et son gap. MAB est une approche basée sur les bandits manchots (*Multi-Armed Bandits*) permettant d'apprendre la meilleure combinaison d'objectifs à optimiser à chaque itération. Nous montrons comment ces solutions peuvent être couplées avec deux modèles courants de simulation d'apprenants : BKT et IRT. Nos expérimentations démontrent la nécessité de prendre en compte les 3 objectifs et d'adapter dynamiquement les objectifs d'optimisation aux capacités de progression de l'apprenant, car MAB permet un taux de maîtrise plus élevé.

ABSTRACT

Multi-objective Test Recommendation for Adaptive Learning

Upskilling is a fast-growing segment of the education economy. Yet, there is little algorithmic work that focuses on crafting dedicated strategies to reach high-skill mastery. In this paper, we formalize ADUP, an iterative upskilling problem that combines mastery learning and Zone of Proximal Development. We design two solutions for ADUP : MOO and MAB. MOO is a multi-objective optimization approach that relies on Hill Climbing to adapt the difficulty of recommended tests to three objectives : learner's predicted performance, aptitude, and skill gap. MAB is a meta approach based on Multi-Armed Bandits to learn the best combination of objectives to optimize at each iteration. We show how these solutions are combined with two common learner simulation models : BKT (KT-IDEM) and Item Response Theory (IRT). Our simulation experiments demonstrate the necessity of leveraging all three objectives and the need to adapt the optimization objectives to the learner's progression ability as MAB offers a higher mastery rate and a better final skill gain than MOO.

MOTS-CLÉS : Apprentissage adaptatif, recommandation, optimisation..

KEYWORDS: Adaptive learning, recommendation, optimization..

ARTICLE : **Accepté à IA-ÉDU@CORIA-TALN 2025.**

1 Contexte

La croissance rapide des nouvelles opportunités d'apprentissage comme les MOOCs, les tutoriels ou les forums communautaires oriente de plus en plus l'attention vers l'amélioration des compétences en ligne. L'*upskilling* (montée en compétences) en dehors des parcours de formation formels constitue aujourd'hui un segment en forte expansion de l'économie de l'éducation. Par ailleurs, les apprenants sont de plus en plus engagés dans un apprentissage auto-dirigé, gérant eux-mêmes de nombreux aspects de leur formation, ce qui implique souvent de travailler de manière autonome sur diverses activités d'apprentissage. Par conséquent, il devient de plus en plus difficile de garantir la qualité des acquis dans ces formats d'apprentissage, car ils peuvent entraîner une compréhension superficielle d'un sujet.

Idéalement, chaque apprenant devrait recevoir des tests choisis de façon à faire progresser réellement ses compétences, en tenant compte de sa capacité à résoudre des exercices en fonction de son niveau et de ses performances passées. C'est précisément l'objectif de l'apprentissage par maîtrise, une stratégie pédagogique qui met l'accent sur le temps nécessaire à chaque apprenant pour acquérir les mêmes compétences et atteindre le même niveau de maîtrise.

2 Contributions

Ce travail publié à *Trans. Large Scale Data Knowl. Centered Syst'24* (Bouarour *et al.*, 2024) prolonge un travail antérieur (Bouarour *et al.*, 2023) en formalisant AdUp (*Adaptive Upskilling*) comme un problème d'optimisation où l'on cherche, à chaque itération, à recommander un ensemble de k -tests à un apprenant. Ces tests doivent maximiser la performance attendue et l'aptitude, tout en minimisant l'écart de compétence accumulé. La combinaison simultanée de ces trois objectifs constitue la nouveauté principale de cette formalisation. Le défi majeur réside dans la nature multi-objectifs du problème. Deux solutions sont explorées :

- MOO (*Multi-Objective Optimization*) : Cette approche repose sur une solution de Pareto basée sur la dominance entre ensembles de tests, et utilise une heuristique de type *Hill Climbing* (Omidvar-Tehrani *et al.*, 2016) pour trouver des solutions non dominées. Diverses variantes peuvent être construites selon les combinaisons d'objectifs. Toutefois, MOO applique les mêmes objectifs tout au long du processus, ce qui limite son adaptabilité.
- MAB (*Multi-Armed Bandits*) : Pour pallier cette rigidité, MAB est introduite comme une approche adaptative qui apprend dynamiquement quels objectifs optimiser à chaque itération. Par exemple, si un apprenant échoue plusieurs fois aux mêmes tests, il serait plus pertinent de privilégier la réduction de l'écart avant de proposer des tests plus difficiles. MAB est ainsi formalisée comme un problème de type bandit manchot, capable d'adapter sa stratégie au comportement de l'apprenant.

3 Expérimentations

Nos expériences visent à évaluer l'efficacité des dimensions d'optimisation sur la montée en compétence (*upskilling*). Pour cela, nous avons divisé notre étude expérimentale en deux parties. Dans la

première partie, nous analysons l'impact de nos solutions sur l'atteinte de la maîtrise en simulant les réponses des apprenants ainsi que l'ensemble du processus d'apprentissage. Nous formulons quatre questions de recherche :

- **RQ1.** La combinaison de toutes les dimensions d'optimisation est-elle bien adaptée pour atteindre la maîtrise et améliorer la progression des compétences ?
- **RQ2.** Les paramètres choisis pour la stratégie de mise à jour des compétences influencent-ils les résultats ?
- **RQ3.** Le choix du modèle de simulation de l'apprenant : Bayesian Knowledge Tracing (BKT) (Corbett & Anderson, 1994) et Item Response Theory (IRT) (Reckase & Reckase, 2009), a-t-il un impact sur la maîtrise et la progression ?
- **RQ4.** L'application d'une méta-stratégie (comme MAB), qui choisit dynamiquement un sous-ensemble de dimensions à optimiser à chaque itération, améliore-t-elle l'atteinte de la maîtrise ?

Les résultats expérimentaux montrent que l'approche MOO permet d'atteindre le plus haut taux de maîtrise en un nombre réduit d'itérations, confirmant ainsi la théorie de la Zone de Développement Proximal (ZPD) et du Flow, ainsi que l'importance d'exploiter l'aptitude pour proposer des défis adaptés aux apprenants. Notre étude révèle que MOO est robuste face aux variations du paramètre N dans la stratégie de mise à jour des compétences (N réponses correctes consécutives). Les conclusions obtenues avec le modèle de simulation BKT sont également généralisables avec IRT, bien que ce dernier favorise la réduction de l'écart, tandis que BKT privilégie la performance attendue.

4 Conclusion

Nous avons abordé la montée en compétence adaptative en suivant une approche d'apprentissage par maîtrise. L'originalité de notre méthode réside dans l'adaptation de la difficulté des tests selon les performances prédites de l'apprenant, son aptitude et ses lacunes. Deux approches ont été proposées : MOO, qui résout directement le problème, et MAB, qui choisit dynamiquement entre différentes variantes d'optimisation à chaque itération. Nos expériences ont montré que MAB permet un meilleur taux de maîtrise et une progression plus importante des compétences finales par rapport à MOO. Cependant, MOO attribue des tests de meilleure qualité et avec plus de précision. Nous avons également testé l'effet de différents modèles de simulation d'apprenants sur la réussite à la maîtrise.

Dans le futur, nous envisageons d'enrichir notre cadre théorique en intégrant d'autres théories de l'apprentissage, telle que l'apprentissage collaboratif, qui a prouvé son efficacité en ligne, notamment via le feedback entre pairs et les discussions entre apprenants. Enfin, nous visons une personnalisation plus fine de l'expérience d'apprentissage en modélisant des profils d'apprenants à partir de leurs performances passées. Ces profils pourraient être utilisés pour attribuer des tests soit par regroupement selon leur niveau général, soit par filtrage collaboratif basé sur les réussites d'apprenants similaires.

Références

BOUAROUR N., BENOURET I. & AMER-YAHIA S. (2024). *Multi-objective Test Recommendation for Adaptive Learning*, In A. HAMEURLAIN, A. M. TJOA, R. AKBARINIA & A. BONIFATI, Édés.,

Transactions on Large-Scale Data- and Knowledge-Centered Systems LVI : Special Issue on Data Management - Principles, Technologies, and Applications, p. 1–36. Springer Berlin Heidelberg : Berlin, Heidelberg. DOI : [10.1007/978-3-662-69603-3_1](https://doi.org/10.1007/978-3-662-69603-3_1).

BOUAROUR N., BENOURET I., D'HAM C. & AMER-YAHIA S. (2023). Adaptive Test Recommendation for Mastery Learning. In *Proceedings of the 2nd International Workshop on Data Systems Education : Bridging Education Practice with Education Research*, DataEd '23, p. 18–23, New York, NY, USA : Association for Computing Machinery. DOI : [10.1145/3596673.3596977](https://doi.org/10.1145/3596673.3596977).

CORBETT A. T. & ANDERSON J. R. (1994). Knowledge tracing : Modeling the acquisition of procedural knowledge. *User modeling and user-adapted interaction*, **4**, 253–278.

OMIDVAR-TEHRANI B., AMER-YAHIA S., DUTOT P.-F. & TRYSTRAM D. (2016). Multi-objective group discovery on the social web. In *Joint European Conference on Machine Learning and Knowledge Discovery in Databases*, p. 296–312 : Springer.

RECKASE M. D. & RECKASE M. D. (2009). Unidimensional item response theory models. *Multidimensional item response theory*, p. 11–55.