

Intelligence Artificielle et Jeux Vidéo

Martin Moritz
Tampere University
Finlande

`martin.moritz@tuni.fi`

Jaakko Peltonen
Tampere University
Finlande

`jaakko.peltonen@tuni.fi`

Les discussions à propos de l'intelligence artificielle (IA) occupent une place grandissante au travail, dans les médias, ou sur les réseaux sociaux. Bien sûr, l'industrie du jeu vidéo, au même titre que les autres, est désormais largement impactée par l'IA générative, celle des modèles de langage (LLMs) ou de diffusion, pour la création de textes, d'images, ou encore de code informatique.

Mais pour les joueurs, l'IA désigne aussi les algorithmes qui leur permettent de se confronter à la machine, et sur lesquels les chercheurs se sont penchés de longue date. D'abord pour les jeux classiques comme les dames, les échecs, ou le Scrabble (Schaeffer, 2001). Ensuite pour les jeux vidéo, où la qualité de l'adversaire gouverné par l'IA est au moins aussi importante que le réalisme des graphiques. Des méthodes traditionnelles telles que l'algorithme A*, ou les modèles d'automates finis, ont fait leurs preuves depuis des décennies. En outre, le jeu vidéo, par sa complexité et sa diversité, offre un terrain idéal au développement de nouvelles méthodes d'IA, répondant aux attentes des joueurs pour différents genres comme les jeux d'actions, d'aventure, ou de stratégie (Laird and VanLent, 2001; Nareyek, 2004).

Contrairement au grand public qui ne s'y est intéressé que relativement récemment, les adeptes de jeux vidéo sont familiers depuis longtemps avec l'IA. Nous avons voulu explorer leurs attitudes, opinions, préoccupations, ou autres idées novatrices. Cette communauté étant active sur les réseaux sociaux, les discussions et débats qui ont trait à l'IA sont une source d'information presque inépuisable.

Notre objectif de recherche se décompose ainsi. 1. Quels sont les principaux thèmes abordés quand les utilisateurs de jeux vidéo discutent de l'IA et de l'IA générative sur la plateforme Reddit ? 2. Pour certains de ces thèmes, peut-on mettre en relief des

opinions convergentes ou divergentes ? 3. Comment la prévalence de ces thèmes varie-t-elle dans le temps ? Est-ce qu'elle s'intensifie, s'atténue, ou connaît des pics ?

Reddit est un forum de discussion généraliste, où un message réplique à un commentaire précédent, ou directement au message original («submission») d'un fil de discussion. Chaque discussion a lieu dans un «subreddit», qui correspond à une communauté d'utilisateurs (Anderson, 2015). Reddit revendique plus de 100 000 communautés alimentées quotidiennement par 130 millions d'utilisateurs¹ qui y discutent de leurs centres d'intérêts. Nous avons collecté, en mars 2026, des commentaires à partir de mots-clefs relatifs au jeu vidéo et à l'IA. Nous avons écarté les publications antérieures à janvier 2021, quelques mois après le lancement de GPT-3 (Brown et al., 2020). De plus nous nous sommes concentrés sur les messages en anglais.

Notre corpus est constitué de 22 866 documents contenant un commentaire répondant directement à une «submission». Chaque document inclut aussi les réponses directes ou indirectes au commentaire initial. Pour révéler les différents sujets de discussion nous avons entraîné un modèle thématique (STM), par Roberts et al. (2014), qui est une amélioration du modèle d'allocation de Dirichlet latente (LDA) par Blei et al. (2003). Un des avantages de STM est la possibilité de prendre en compte des métadonnées lors de l'apprentissage. Nous utilisons pour cela la date de publication du commentaire, ainsi qu'un partitionnement des auteurs basé sur le texte des messages qu'ils ont récemment postés sur Reddit. Nous avons ainsi extrait six catégories («Developer», «Gamer», «Hardware», «Leisure», «LLMs», et «Politics»).

1. <https://redditinc.com/press>

STM et LDA sont des *modèles génératifs probabilistes*. Ils supposent que les documents du corpus sont générés par le processus aléatoire qui suit. Pour chaque mot d'un document, un thème est tiré au sort, selon une distribution propre à ce document. Puis un mot est tiré au sort selon une distribution propre à ce thème. L'apprentissage consiste à déterminer des valeurs pour ces deux familles de distributions qui ont le plus de chance de générer le corpus de documents observé. Il est bon de mentionner que ces modèles ne sont capables de générer un échantillon du texte observé, ni même un quelconque texte qui fasse sens, qu'avec une probabilité infinitésimale. Ils diffèrent donc fondamentalement des LLMs puisque l'intérêt ici réside dans les valeurs des paramètres appris, et non dans la qualité des textes que le modèle peut prédire. Dans notre cas, ces paramètres nous permettent, entre autres, d'interpréter des thèmes émanant de notre corpus de document à partir de leurs mots proéminents, ainsi que de mesurer l'évolution dans le temps de la proportion avec laquelle ils sont discutés.

L'entraînement d'un modèle thématique n'est pas déterministe. Dans la plupart des cas on doit se poser au préalable la question du nombre de thèmes à rechercher. Et, après entraînement d'un échantillon de modèles, il faut trouver un moyen d'en sélectionner un. Dans le cas de données réelles, il n'y a tout bonnement pas de réponses toutes faites à ces deux problèmes. Ici, nous utilisons une des fonctionnalités de la version R du package STM décrite dans [Roberts et al. \(2019\)](#). Cette méthode sélectionne automatiquement le nombre de thèmes. En entraînant 50 modèles, leur nombre a varié de 72 à 105. Puis nous avons sélectionné le modèle offrant la meilleure cohérence thématique ([Röder et al., 2015](#)) qui comporte 95 thèmes. L'analyse des meilleurs mots ainsi que la lecture des documents les plus représentatifs de chaque thème nous permet de leur donner un titre interprétatif. Neuf de ces thèmes présentaient peu d'intérêt pour notre étude, et nous avons regroupé les 86 restants en 5 catégories. 1. L'IA comme adversaire du joueur 2. L'IA pour générer des ressources graphiques et audio 3. L'IA pour générer du texte naturel ou des mondes ouverts 4. L'IA pour générer du code et pour le développement de jeux 5. Impact de l'IA générative sur le marché du jeu vidéo.

Ces discussions révèlent de nombreux débats

liés à l'IA et au jeu vidéo. Peut-on tolérer son usage par des hobbyistes et le réprouver pour les grosses productions ? Quel est son bénéfice par rapport aux ressources graphiques disponibles en ligne depuis longtemps ? Étant donné que l'IA est perçue négativement, en raison de son impact sur l'environnement ou sur l'emploi, n'aurait-on pas intérêt à dissimuler son usage ?

Références

- Katie Elson Anderson. 2015. Ask me anything : what is Reddit ? *Library Hi Tech News* 32(5) :8–11.
- David M Blei, Andrew Y Ng, and Michael I Jordan. 2003. Latent Dirichlet allocation. *Journal of Machine Learning Research* 3(Jan) :993–1022.
- Tom Brown, Benjamin Mann, Nick Ryder, Melanie Subbiah, Jared D Kaplan, Prafulla Dhariwal, Arvind Neelakantan, Pranav Shyam, Girish Sastry, Amanda Askell, et al. 2020. Language models are few-shot learners. *Advances in Neural Information Processing Systems* 33 :1877–1901.
- John Laird and Michael VanLent. 2001. Human-level AI's killer application : Interactive computer games. *AI Magazine* 22(2) :15–15.
- Alexander Nareyek. 2004. AI in computer games : Smarter games are making for a better user experience. What does the future hold ? *Queue* 1(10) :58–65.
- Margaret E Roberts, Brandon M Stewart, and Dustin Tingley. 2019. stm : An R package for structural topic models. *Journal of Statistical Software* 91 :1–40.
- Margaret E Roberts, Brandon M Stewart, Dustin Tingley, Christopher Lucas, Jetson Leder-Luis, Shana Kushner Gadarian, Bethany Albertson, and David G Rand. 2014. Structural topic models for open-ended survey responses. *American Journal of Political Science* 58(4) :1064–1082.
- Michael Röder, Andreas Both, and Alexander Hinneburg. 2015. Exploring the space of topic coherence measures. In *Proceedings of the Eighth ACM International Conference on Web Search and Data Mining*. pages 399–408.
- Jonathan Schaeffer. 2001. A gamut of games. *AI Magazine* 22(3) :29–29.