



Fiche pédagogique

Activité : Cadavre exquis 2.0. Le jeu du mot suivant.

Objectifs pédagogiques :

- Comprendre le principe général de fonctionnement d'une IA générative de texte comme ChatGPT.
- Découvrir l'importance du **contexte** (le prompt) dans la génération d'un texte.
- Expérimenter la construction collective d'une réponse, un mot à la fois.
- Comprendre qu'une même question peut conduire à plusieurs réponses plausibles.
- Développer les capacités d'écoute, de coopération et d'anticipation.
- Notion générale de probabilité.

Notions abordées : IA générative, prompt, contexte, probabilité.

Matériel nécessaire : Accès à un logiciel d'IA générative, éventuellement écran ou vidéo-projecteur pour projeter.

Niveau : À partir du cycle 3.

Durée : 30 à 45 minutes.

Messages à faire passer

- Une IA générative de texte comme ChatGPT ne rédige pas toute une phrase d'un seul coup : elle génère le texte mot après mot.
- Pour choisir chaque nouveau mot, elle s'appuie sur le prompt et sur tous les mots déjà générés.
- Plus le contexte est riche et précis, plus la réponse a de chances d'être pertinente.
- Une IA ne produit pas toujours exactement la même réponse : plusieurs suites de mots peuvent être cohérentes.

Déroulement : Cette activité se déroule en **deux manches** afin de montrer l'importance du contexte dans la génération d'un texte.

Etape 1 : Introduction (~ 5 min)

Présenter l'objectif : : « *Aujourd'hui, nous allons essayer de faire comme une intelligence artificielle. Ensemble, nous allons fabriquer une réponse en ajoutant un seul mot chacun, exactement comme le fait une IA générative produit un texte.* »

Expliquer ensuite qu'une IA ne connaît pas la phrase complète à l'avance. Elle construit sa réponse progressivement. À chaque étape, elle choisit un mot qui lui semble cohérent avec le prompt et avec tous les mots déjà écrits.

Mentionner que les IA sont en constante évolution. Elles sont capables de générer du texte, mais aussi de citer les sources, fournir le lien de la page internet et sont aussi connectées à des outils externes pour produire des réponses plus sophistiquées.



Etape 2 : Première manche : le prompt caché (~ 10 min)

Le médiateur écrit une question sur une ardoise, mais ne la montre qu'aux deux premiers participants.

Exemples : Que ferais-tu si tu avais une baguette magique ? Quel est ton animal préféré ? Pourquoi les abeilles sont-elles importantes ? Comment préparer un gâteau au chocolat ? Que ferais-tu si tu pouvais voyager dans le temps ?

Les deux premiers participants connaissent la question. Tous les autres ne voient que les mots proposés au fur et à mesure. Les participants construisent collectivement une réponse :

- chaque participant ajoute un seul mot ;
- il est interdit de modifier les mots précédents ;
- chacun doit proposer le mot qui lui semble le plus logique au vu de ce qu'il connaît.

À la fin, le médiateur révèle la question.

S'ensuit une discussion où le médiateur interroge les participants avec les questions suivantes :

- La réponse répond-elle réellement à la question ?
- À quel moment est-elle devenue cohérente (ou incohérente) ?
- Quels participants avaient le plus d'informations ?
- Pourquoi les derniers participants avaient-ils plus de facilité que les premiers ?

Cette dernière question fait déjà apparaître la notion de contexte.

Etape 3 : Deuxième manche : le prompt partagé (~ 10 min)

Le médiateur propose une nouvelle question. Cette fois-ci, tous les participants voient le prompt dès le départ. Le déroulement est identique : un mot par participant ; aucune modification possible ; chacun essaie de proposer le mot le plus cohérent.

Comparer ensuite les deux productions.

Questions possibles :

- Quelle réponse est la plus cohérente ?
- Pourquoi ?
- Est-il plus facile de proposer un mot lorsque tout le monde connaît la question ?
- Qu'a changé le fait de partager le prompt ?

Conclure : Ce n'est pas parce que les participants sont meilleurs dans la deuxième manche : c'est parce qu'ils disposent de davantage de contexte. C'est exactement ce qui se passe avec une IA générative.

Etape 4 : Conclusion (~ 5 min)

Faire le lien avec les modèles de langage. Expliquer que ChatGPT fonctionne selon un principe similaire :

- il reçoit un prompt ;
- il lit tous les mots déjà présents dans la conversation ;
- il estime quels mots seraient les plus cohérents pour poursuivre le texte ;
- il choisit l'un d'entre eux puis recommence jusqu'à produire toute la réponse.

Le premier exercice montre ce qui se passe lorsque le contexte est incomplet. Le second montre qu'un prompt partagé et plus riche permet de produire une réponse plus pertinente.

Etape 5 : Démonstration avec une IA Pour terminer, le médiateur peut refaire l'exercice avec une IA générative de son choix.

Exemple de prompt : Réponds à la question suivante comme si tu étais un participant de CM2. Ta réponse doit contenir exactement X mots. (X correspondant au nombre de participants.)

Le médiateur peut ensuite comparer la réponse construite par les participants et la réponse produite par l'IA. Cette comparaison permet d'observer que, comme les participants, l'IA construit une réponse cohérente **à partir du contexte fourni**, mot après mot.

Conseils d'animation : Choisir des questions ouvertes qui admettent de nombreuses réponses possibles. L'objectif n'est pas de trouver "la bonne réponse", mais de produire une réponse cohérente collectivement.

Variante : Le contexte caché. Reprendre la réponse fournie à un des exercices ci-dessus, mais en supprimant les premiers mots. Par exemple : Si à la question "Quel est ton animal favori ?", la réponse construite était "Le cheval parce qu'il est facile de jouer avec lui.", on raccourcit en "il est facile de jouer avec lui." Demander de quoi on parle. Les participants constateront qu'il est difficile de comprendre sans le début de la phrase. Puis réafficher toute la phrase. Cette variante illustre très concrètement qu'un modèle de langage s'appuie en permanence sur tout le contexte disponible pour interpréter et générer un texte.