

ÉVALUATION AUTOMATIQUE DU RETOUR À LA SOURCE DANS UN CONTEXTE HISTORIQUE

Les débats parlementaires de la Troisième République française

Aurélien Pellet^{1 2}, Julien Perez², Marie Puren^{2 3}

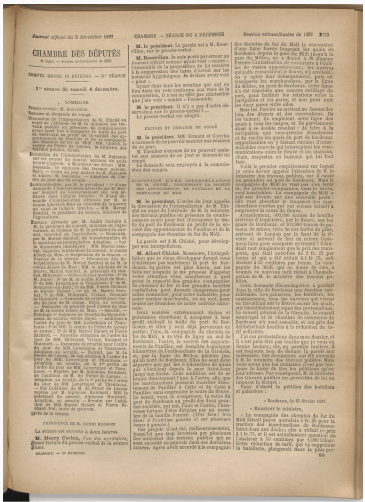
4 juillet 2025 - Journée HNIA à PFIA

¹Epitech ²LRE, Epita ³Centre Jean-Mabillon, Ecole nationale des chartes

LES DÉBATS PARLEMENTAIRES DURANT LA III^E RÉPUBLIQUE

Compréhension des motivations des députés, des oppositions aux projets de loi, et des sujets traités

- Analyse thématique du corpus et exploitation des données sérielles
- Interrogation des documents pour répondre à des questions spécifiques
- Annotation automatique pour faciliter l'exploration du corpus
- Validation des affirmations et données extraites en revenant aux sources



Séance parlementaire du 4 décembre 1897

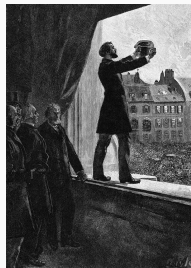
LES LLMS PEUVENT-ILS NOUS AIDER ?

Contexte

- Troisième République : archives parlementaires riches mais volumineuses
- Test des LLM pour l'exploration de corpus historiques

Questions de recherche

- Comment les LLM aident-ils l'exploration d'archives ?
- Comment évaluer la pertinence et la frugalité ?
- Mesurer le retour à la source



Gambetta proclamant la République (Frontispice de Mary Platt Parmele, A Short History of France, 1907)



Séance du 4 Décembre 1897 :

« Il n'y a pas d'affaire Dreyfus [...] il n'y a pas en ce moment et il ne peut pas y avoir d'affaire Dreyfus. »

« l'anarchisme.[...]a ses deux versants : sur l'un, la doctrine fleurit dans toute sa pureté mystique; sur l'autre, l'action apparaît avec toute la brutalité révolutionnaire.[...]C'est aussi sur les théoriciens que vous avez déchaîné votre police, et votre erreur est d'autant plus condamnable, monsieur le ministre, que le régime sous lequel vous la pratiquez est issu lui-même de la double gestation de la violence et de la pensée.

Il y a du sang sur toutes les doctrines, même sur celles que vous invoquez contre le sang! »

— M. Clovis Hugues

Extrait d'un discours sur les répressions contre les Anarchistes
à la Chambre des députés (27 janvier 1894).

UN EXEMPLE D'OPINION TRANCHÉE

« Aujourd'hui je demande à un paysan :

— Qu'est-ce que la patrie ? Est-ce ton champ ?

— Il répondra : Oui.

Cependant, il sait bien qu'il fait partie d'une nation qui s'appelle la France ; mais comment la connaît-il ?

Un jour, on publie à son de trompe l'appel des conscrits. Son fils va prendre un numéro dans une boîte. Soldat pendant cinq ans, c'est la patrie qui le demande.

De temps en temps, il voit apparaître un gendarme qui lui dresse un procès-verbal : c'est une nouvelle incarnation de la patrie.

Les charges, il les conçoit très bien. Quant aux avantages, il les aperçoit moins distinctement.

L'idée de la patrie est fort obscurcie chez lui, et quand, au moment d'une guerre, on vient lui parler en phrases pompeuses de ses obligations, il répond :

« Mes obligations, je les connais ; mais quelles sont les obligations de la patrie envers moi ? »

— Yves Guyot, Les Droits de l'Homme, 27 mars 1876.

Cité par Clovis Hugues, discours à la Chambre des députés, 27 janvier 1894.

Un corpus « massif » : 15 législatures entre 1881 et 1940 à la Chambre des députés, avec 10 000 à 12 000 images par législature

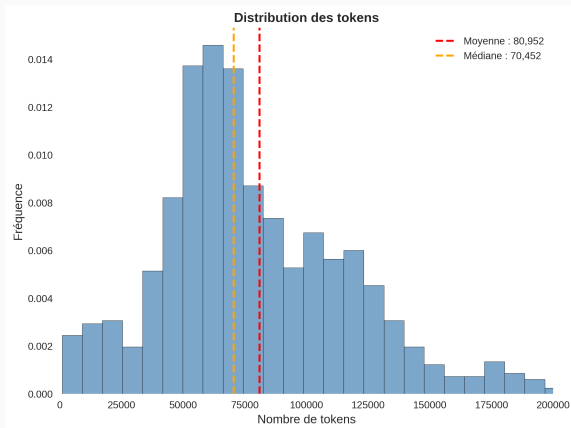
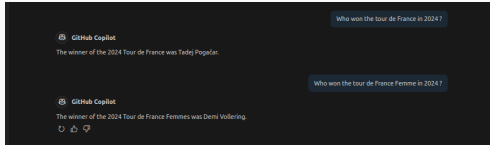


Figure – Distribution du nombre de tokens

LE RAG POUR LA RECHERCHE HISTORIQUE



Réponses erronées de GitHub Copilot

Pour pallier ces limites : agentification des LLM

- Recherche sur Internet
- Interrogation d'API
- RAG (Retrieval-Augmented Generation)

Limites

- Modèles fermés
- Limite dans le temps
- Encodage du savoir



Résultat correct via recherche web

LE PROBLÈME DE LA SPÉCIFICITÉ ET DE LA TAILLE

- Les architectures Transformer sont limitées par la taille maximale de la séquence qu'elles peuvent traiter.
- Le risque d'**hallucination** apparaît bien avant d'atteindre cette limite.

Modèle	Contexte max.	Date de sortie
BERT	512 tokens	Octobre 2018
GPT-3.5	4 096 tokens	Mars 2023
LLaMA 2	4 096 tokens	Juillet 2023
LLaMA 3.1 (70B)	128 000 tokens	Juillet 2024
GPT-4o	128 000 tokens	Mai 2024
Qwen2.5	128 000 tokens	Septembre 2024
Gemma-3 :27b	128 000 tokens	Mars 2025

Capacités contextuelles des modèles de langue

LE PROBLÈME DE LA SPÉCIFICITÉ ET DE LA TAILLE

Query	Model	Response
Quelles sont les deux propositions de résolutions déposées par M. Louis Marin ?	GPT-3.5 Turbo	Code d'erreur : 400 - {'error': {'message': "La longueur maximale du contexte de ce modèle est de 16385 tokens. Cependant, vos messages ont généré 49791 tokens."}}
	LLaMA 3.1 70B	Les deux propositions de résolution déposées par M. Louis-Marin concernent la pétition n° 1615 et portent sur le renvoi de cette affaire à la commission des réclamations, avec un vote différent pour chacune : l'une est favorable à un renvoi simple (proposition n° 1) et l'autre préconise un renvoi avec une déclaration formelle (proposition n° 2).
Qui est le président de la séance du jour ?	LLaMA 3.1 70B (full)	M. André
	LLaMA 3.1 70B (5k chars)	M. Jules Méline

Limites des LLM sur les séances des débats parlementaires

Points clés

- **Contexte limité** : Les modèles comme GPT-3.5 ne disposent pas d'une taille de contexte suffisante pour traiter un débat complet.
- **Risque d'hallucination** : Même avec un modèle très récent et une taille de contexte plus grande, le risque d'hallucination reste élevé.

RETRIEVAL-AUGMENTED GENERATION (RAG)

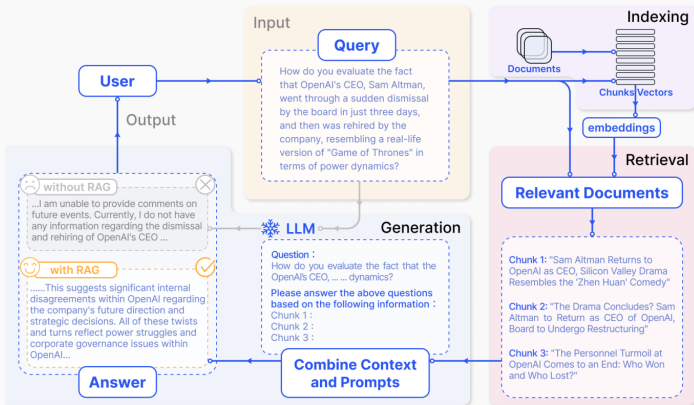


Schéma d'un RAG simple, Yufan Gao, Retrieval-Augmented Generation for Large Language Models : A Survey, 2023

EXEMPLE DE RAG NAÏF

1. Question

Quel est le montant total des dégâts subis par les routes locales du département des Côtes-du-Nord ?

3. Template de Prompt

SYSTÈME : Je vais vous fournir des extraits de débats parlementaires de la Troisième République. Analysez-les et utilisez-les pour répondre à la question donnée.

UTILISATEUR :

- Extrait 1 : [Premier chunk]
- ...
- Extrait K : [K-ème chunk]

Question : [Question]

2. Phase de Récupération

Modèle d'embedding : mGTE

Documents récupérés :

[**Chunk n°1**] Les dégâts sont beaucoup plus considérables; d'après un document officiel — le rapport du préfet à la commission départementale — la dépense pour les routes locales à la charge du département s'élève à 275 060 francs.

[...]

[**Chunk n°k**] ...

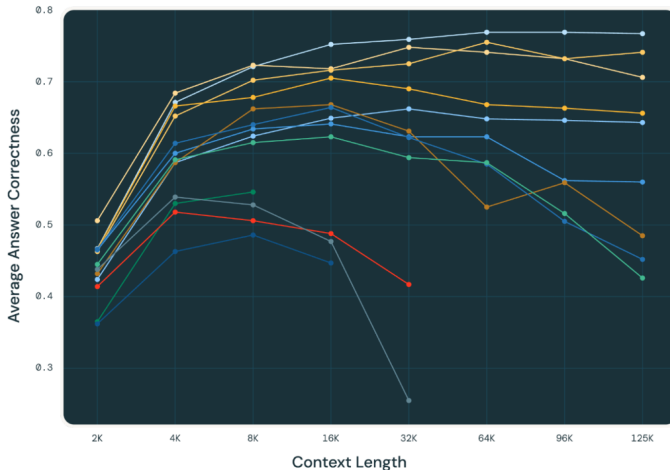
4. Réponse Générée

Le montant total des dégâts subis par les routes locales du département des Côtes-du-Nord est de **275 060 francs**.

EXEMPLE DE QUESTIONS

Source	Query	Chunk retrieved (top 1 chunk)	Answer
1881-01-22	Quel est le montant total des dégâts subis par les chemins vicinaux dans le département des Côtes-du-Nord ?	Les dégâts sont bien plus considérables; ils s'élèvent, — et j'ai un document officiel, c'est le rapport du préfet à la commission départementale, — rien que pour les chemins vicinaux qui sont à la charge du département, ils s'élèvent à la somme de 275,060 fr.	Le montant total des dégâts subis par les chemins vicinaux dans le département des Côtes-du-Nord s'élève à 275 000 francs.
1881-01-22	"Quel est le montant de l'emprunt autorisé pour le département de l'Allier et à quoi sera-t-il utilisé?"	t'rPROJET«Art. 1er. — Le département de l'Allier est autorisé, conformément la demande que le conseil général en a faite, à emprunter la caisse des lycées, collèges et écoles, aux conditions de cet établissement, une somme de 310,000 fr. applicable au paiement de subventions destinées aux travaux des lycées de Moulins et de Montluçon.	Le montant de l'emprunt autorisé pour le département de l'Allier est de 310,000 fr. Il sera utilisé pour le paiement de subventions destinées aux travaux des lycées de Moulins et de Montluçon.
1882-03-23	"Quel est le montant de l'emprunt autorisé pour le département des Bouches-du-Rhône?"	1882-03-23	Le montant de l'emprunt autorisé pour le département des Bouches-du-Rhône n'est pas mentionné dans les extraits de débat parlementaire fournis.
	Qui est le président de la séance du jour?	PRÉSIDENCE DE M. JULES MÉLINE La séance est ouverte à deux heures un quart. M. Charles Chevalier (Manche) , l'un des secrétaires, donne lecture du procès-verbal de la séance d'hier.	M. Jules Méline

LLM, CONTEXTE ET HALLUCINATION



Long

Context RAG Performance of LLMs, databricks,
<https://www.databricks.com/blog/long-context-rag-performance-llms>

Inférence et frugalité

- Limites des LLM généralistes : inefficaces sur des corpus historiques longs et spécialisés.
- Objectif du RAG : maximiser la qualité des réponses tout en limitant la taille des entrées.

Objectifs de l'étude

- **Chunking** : le découpage du texte influence la qualité du retrieval et des réponses.
- **Évaluation** : comparaison de plusieurs stratégies de découpage dans un pipeline RAG.
- **Reranking** : étude de l'impact du reranker sur le retour à la source.

ÉVALUATION DES CAPACITÉS DE RETRIEVING

GÉNÉRATION AUTOMATIQUE DE QUESTIONS

Pour explorer les capacités de retrieving sur de longs documents et constituer un corpus conséquent de questions, nous utilisons un LLM pour générer des questions.

- À partir de débats prédécoupés (pour réduire leur taille), un LLM est utilisé pour **générer une liste de questions** et **fournir le passage du débat utilisé** pour chaque question.
- Llama 3.1 :70B est utilisé pour la génération de questions.
- Vérification systématique que la source renvoyée par le LLM est bien présente dans le texte d'origine.
- Le prompt inclut une phase de réflexion et de vérification pour améliorer la qualité des résultats.
- On cherchera à réduire les biais potentiels dans la génération en **retravaillant le prompt et en validant les questions générées**.

REFORMULATION ET ENRICHISSEMENT DES QUESTIONS

Question initiale	Reformulation	Nouvelles questions	Reformulation
Selon le comte de Douville-Maillefeu, quel facteur est primordial pour assurer la victoire d'une nation dans un conflit, même face à un adversaire de force égale ?	En 1881, lors des discussions sur le budget de la marine, quel est le facteur primordial, selon le comte de Douville-Maillefeu, pour assurer la victoire d'une nation dans un conflit, même face à un adversaire de force égale ?	- Quel arsenal le comte de Douville-Maillefeu propose t-il de supprimer lors des discussions sur le budget de la marine en 1881 ? - Quels sont les investissements que le comte de Douville-Maillefeu propose de revoir lors des discussions sur le budget de la marine en 1881 ?	Ajout de la date et du sujet de la discussion.

Exemple de reformulation et génération manuelle de sous-questions

- Pour le moment, l'évaluation est réalisée par une seule experte du domaine. Il paraît nécessaire de **réaliser l'évaluation à plusieurs mains** et de croiser les résultats.
- Les questions générées se concentrent sur des **éléments factuels, faciles à vérifier et non controversés**.
- Cette méthode ne remplace pas la génération manuelle de questions, qui demeure **le moyen le plus fiable** de produire un corpus de qualité.

La question reste : comment évaluer **les stratégies de chunking dans leur capacité à fournir des informations suffisantes** (retrieving) pour répondre correctement à une question ?

Questions : 1400 questions « single-hop » issues de débats parlementaires, avec traçabilité de leur source d'origine.

Base vectorielle ChromaDB (L2) – **LLM** : Gemma 3 - 27B – **Reranker** : BGE-v2-M3 – **Embeddings** : mGTE – **GPU** : NVIDIA A40 (46 Go)

Métriques :

$$\text{Recall@k} = \frac{|\text{Rel} \cap \text{Ret@k}|}{|\text{Ret@k}|} \quad | \quad \text{nDCG@k} = \frac{\text{DCG@k}}{\text{IDCG@k}}, \text{ où } \text{DCG@k} = \sum_{i=1}^k \frac{\text{rel}_i}{\log_2(i+1)}$$

Méthode	Description	Implémentation	Taille moyenne des chunks (nombre de tokens)
S	Découpage par sections	Une Regex identifie les débuts de sections par séquences de majuscules	$10\,989 \pm 14\,302$
S&PP	Sections + Prises de parole	Prises de parole identifiées par le préfixe « M. »	808 ± 1267
S&PP&MaxN	Sections + Prises de parole + Limite de taille	Découpage interrompu si le chunk contient moins de N caractères	$204 \pm 66 - 1412 \pm 940$ (N = 1000 – N = 10 000)

Description des stratégies de chunking étudiées. Pour chaque méthode, un chunking récursif est appliqué, et les frontières entre les différentes sections sont identifiées à l'aide d'expressions régulières.

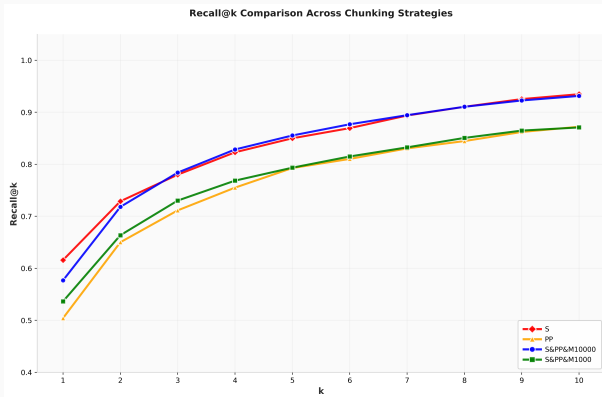
DÉTAIL DU CHUNKING

Extrait d'un chunk en fonction de différentes stratégies	S	S&PP	S&PP&Max10000
	PRÉSIDENCE DE M. HENRI BRISSON La séance est ouverte à deux heures. M. Henry Cochin, l'un des secrétaires, donne lecture du procès-verbal de la séance d'hier. M. le président. La parole est à M. Bourrillon, sur le procès-verbal. M. Bourrillon. Je suis porté par erreur au Journal officiel comme ayant voté « contre » l'ensemble de la proposition de loi relative à la suppression des taxes d'octroi sur les boissons hygiéniques. Je déclare avoir voté « pour ». Ayant dans tous les scrutins qui ont eu lieu émis un vote conforme aux vues de la commission, il n'est en effet pas admissible que j'aie voté « contre » l'ensemble. M. le président. Il n'y a pas d'autre observation sur le procès-verbal? Le procès-verbal est adopté.	Idem colonne S.	PRÉSIDENCE DE M. HENRI BRISSON La séance est ouverte à deux heures. M. Henry Cochin, l'un des secrétaires, donne lecture du procès-verbal de la séance d'hier.
Nombre de chunks	59	79	255

Impact des stratégies sur la taille des chunks

RÉSULTATS

STRATÉGIES DE CHUNKING VS RECALL@K



Chunking vs Recall@k

Points clés

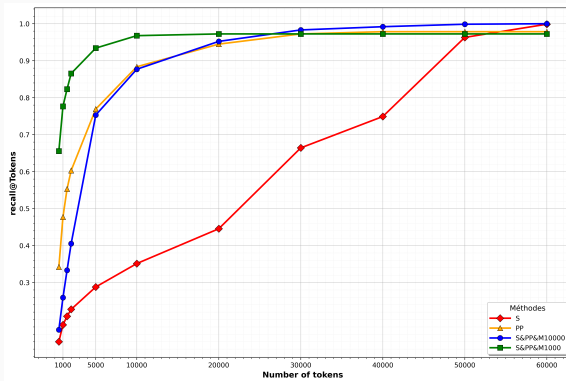
- S offre les meilleurs résultats grâce à des chunks très longs (11k tokens), mais limite le retour à la source.
- **S&PP&Max10000** constitue un bon compromis : deuxième meilleure performance avec bien moins de tokens, et un recall proche de S dès 3 documents.
- **S&PP&Max1000** surpasse **PP** malgré des chunks 4× plus petits, probablement grâce à une taille plus homogène favorisant des embeddings de meilleure qualité.
- Pour les petits modèles, **S&PP&Max1000** semble la stratégie la plus adaptée.

DÉTAIL DES RÉSULTATS

Stratégie de Chunking		Taille moyenne des chunks (nombre de tokens)	Recall@k					nDCG@k	
Granularité	Taille limite (caractères)		1	2	3	5	10	5	10
S&PP	1000	204 ± 66	0.54	0.66	0.73	0.79	0.87	0.68	0.70
	10000	1412 ± 940	0.58	0.72	0.78	0.86	0.93	0.73	0.74
S PP	None	10989 ± 14302	0.62	0.73	0.78	0.85	0.93	0.74	0.77
		808 ± 1267	0.50	0.65	0.71	0.79	0.87	0.66	0.68

Évaluation du retriever pour différentes stratégies de chunking. Pour chaque configuration, nous calculons le nombre moyen de tokens par chunk récupéré sur l'ensemble des questions. La stratégie de segmentation par section donne presque toujours les meilleurs résultats, mais elle est très coûteuse en nombre de tokens. La stratégie **S&PP&Max10000** offre un bon compromis entre coût à l'entrée et performance du retriever.

STRATÉGIES DE CHUNKING VS RECALL@TOKENS

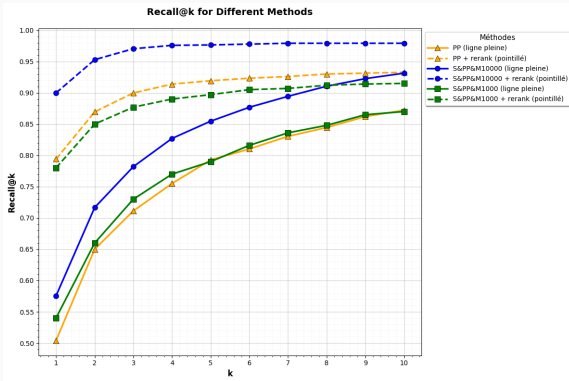


Évaluation des stratégies de chunking,
incluant les questions sans documents récupérés.

Points clés

- Évaluation à nombre de tokens constant : pour chaque stratégie, on récupère un maximum de documents sans dépasser un seuil de tokens fixé. Une question pour laquelle aucun chunk n'est récupéré est considérée comme un échec du retriever.
- S est désavantagée : ses chunks très longs diluent l'information et ne s'adaptent pas aux contextes de petite taille.
- S&PP&Max1000 apparaît comme la stratégie la plus adaptée pour un faible nombre de tokens en entrée. À partir de 5000 tokens, S&PP&Max10000 et PP offrent des performances similaires, mais la plus faible variance des chunks générés par la première peut en faire une option plus robuste.

L'APPORT DU RERANKER



Comparaison des résultats sur le retriever avec et sans reranking. En ligne pleine les résultats avec un RAG classique et en pointillé avec un RAG augmenté par reranking

Points clés

- Pour chaque question, 20 documents sont récupérés, puis reclassés avec le modèle **bge-reranker-v2-m3** selon un score de similarité (question, document). Le recall est recalculé après reclassement.
- Le reranking améliore nettement les performances : pour $k = 1$, le recall de **PP** passe de 0.5 à 0.79. À partir de $k = 3$, cette stratégie gagne en moyenne 2 points sur **S&PP&Max1000**.
- Le reranker corrige les imprécisions des embeddings, rendant **PP** plus performante. **S&PP&Max10000** reste compétitive, avec un recall supérieur d'au moins 5 points sur l'ensemble des évaluations.

DÉTAIL DES RÉSULTATS

Stratégie de Chunking		Taille moyenne des chunks (nombre de tokens)	Coût de calcul (temps)	Recall@k			
Granularité	Taille limite (caractères)			1	2	5	10
S	—	10989 ± 14302	<1min	0.62	0.73	0.85	0.93
S&PP	1000	204 ± 66	<1min	0.54	0.66	0.79	0.87
	10000	1412 ± 940	<1min	0.58	0.72	0.86	0.93
PP	—	808 ± 1267	<1min	0.50	0.65	0.79	0.87
S&PP+reranking	1000	204 ± 66	9min	0.78	0.85	0.89	0.92
	10000	1412 ± 940	1h47min	0.90	0.95	0.98	0.98
PP+reranking	—	808 ± 1267	2h25min	0.79	0.87	0.92	0.93

Évaluation du retriever pour différentes stratégies de chunking. Pour chaque configuration, nous calculons le nombre moyen de tokens par chunk récupéré sur l'ensemble des questions. La stratégie de segmentation par section donne presque toujours les meilleurs résultats, mais elle est très coûteuse en nombre de tokens. La stratégie S&PP&Max10000 offre un bon compromis entre coût à l'entrée et performance du retriever.

Points clés

- Le coût du reranking est élevé : environ 2h d'inférence pour **S&PP&Max10000** et **PP**. **S&PP&Max1000** constitue un bon compromis : pas de reranking, temps d'inférence < 10 min, et résultats proches des meilleures stratégies.

- Améliorer la phase de génération de questions
- Produire des questions portant sur plusieurs sections d'un même débat ou sur plusieurs débats
- Après la phase de retrieval, évaluer la qualité de la réponse générée dans le cadre du RAG
- Poursuivre l'annotation automatique et introduire de nouvelles tâches pour comparer les stratégies de chunking

Aurélien Pellet : `aurelien.pellet@epita.fr`

Julien Perez : `julien.perez@epita.fr`

Marie Puren : `marie.puren@epita.fr`