

LE BONHEUR EST SUR TWITTER

Un baromètre du moral des Français

LE BONHEUR EST SUR TWITTER

Un baromètre du moral des français

THOMAS RENAULT

RUED'ULM

Nous appliquons dans ce livre la plupart des rectifications orthographiques
de la dernière réforme de l'Académie (JO du 6 décembre 1990).

© Éditions Rue d'Ulm/Presses de l'École normale supérieure, 2022
45, rue d'Ulm – 75230 Paris cedex 05
www.pressens.fr
ISBN 978-2-7288-0782-6
ISSN 1951-7637

Le Cepremap est, depuis le 1^{er} janvier 2005, le CEntre Pour la Recherche EconoMique et ses APplications. Il est placé sous la tutelle du ministère de la Recherche. La mission prévue dans ses statuts est d'assurer *une interface entre le monde académique et les décideurs publics*.

Ses priorités sont définies en collaboration avec ses partenaires institutionnels : la Banque de France, le CNRS, France Stratégie, la direction générale du Trésor, l'École normale supérieure, l'INSEE, l'Agence française du développement, le Conseil d'analyse économique, le ministère chargé du Travail (DARES), le ministère chargé de l'Environnement, (ADEME), le ministère chargé de la Santé (DREES) et la direction de la recherche et de l'innovation du ministère de la Recherche.

Les activités du Cepremap sont réparties en trois axes. Axe 1 : Macro-économie ; Axe 2 : Bien-être, travail et politique publique ; Axe 3 : Mondialisation, développement et environnement. Chaque axe dispose d'un observatoire qui coordonne les recherches menées en propre par le Cepremap, et des programmes de recherche qui regroupent une centaine de chercheurs, cooptés par les animateurs des programmes, au sein notamment de l'École d'économie de Paris.

La coordination de l'ensemble des axes est assurée par *Claudia Senik*.

L'affichage sur Internet des documents de travail réalisés par les chercheurs dans le cadre de leur collaboration au sein du Cepremap tout comme les opuscules publiés en collaboration avec les éditions Rue d'Ulm visent à rendre accessible les recherches portant sur la politique économique.

Daniel COHEN
Directeur du Cepremap

Sommaire

Introduction	11
<i>Vox populi</i>	11
1. Comment mesurer le moral des Français en ligne ?	17
<i>Construire une base de données de tweets</i>	18
<i>Mesurer les sentiments et les émotions</i>	22
2. Le moral des Français au prisme de Twitter	27
<i>Le moral global des Français sur Twitter : du mois à la journée</i>	27
<i>Sentiments et émotions</i>	31
<i>Les heures et les jours</i>	32
<i>Genre, politique et émotions</i>	34
<i>Les femmes dans la pandémie</i>	35
<i>Le centre et les extrêmes : la polarisation politique des émotions</i> ..	37
<i>Polarisation de l'espace médiatique</i>	41
3. Robustesse et extensions de l'indicateur	49
<i>Comparaison avec l'indice Hedonometer</i>	49
<i>Indicateur de sentiment et indicateurs de bien-être</i>	50
<i>Que mesure-t-on en analysant le contenu des messages ?</i>	51
<i>Localisation et origine des tweets</i>	54
4. Pour ceux qui ont (encore) des doutes sur la validité des données Twitter	57
<i>Mesurer le bien-être via Twitter</i>	57
<i>Les limites des données en ligne</i>	62
Conclusion	67
Annexe 1. Analyse de sentiment	69

Annexe 2. La mesure du bien-être <i>via</i> Google et Facebook	79
Annexe 3. L'impact négatif causal de l'utilisation des réseaux sociaux sur le bien-être	81
Liste des figures, tableaux et encadré	83
Bibliographie	85

EN BREF

Cet ouvrage présente le baromètre du moral des Français élaboré pour l'Observatoire du bien-être du Cepremap sur la base des messages échangés sur Twitter. Sur les séries ainsi élaborées, on peut lire la dégradation du score de sentiment à partir de l'automne 2018, prélude à la crise des Gilets jaunes. En classant les internautes selon les médias et les personnalités politiques auxquels ils sont abonnés, on constate une polarisation multiple de l'espace des émotions : d'une part, le contraste politique bien identifié entre le centre et les extrêmes, ces derniers étant caractérisés par un bien-être plus faible et des émotions plus négatives et, d'autre part, un contraste générationnel, avec des médias jeunes associés à une plus forte expression de sentiments négatifs. Ce mal-être a sans doute trouvé une traduction dans les urnes au premier tour de l'élection présidentielle de 2022. En permettant de quantifier la tonalité des émotions exprimées sur Twitter, ces indicateurs proposent un nouveau regard sur les facteurs influençant le moral des Français. Ils sont désormais en ligne et à la disposition du public sur le site de l'Observatoire du bien-être.

Thomas Renault est maître de conférences à l'Université Paris I Panthéon-Sorbonne et conseiller scientifique au Conseil d'analyse économique. Ses recherches portent principalement sur l'analyse des réseaux sociaux et l'utilisation de méthodes d'analyse textuelle et de machine learning pour la prévision en économie et en finance.

Remerciements : Cette étude a bénéficié du travail de Romain Jouhameau – assistant de recherche au Cepremap – ainsi que de la relecture et de la rédaction de Mathieu Perona.

Introduction

VOX POPULI

En France, plus de 12 millions de personnes se connectent chaque mois sur Twitter pour s'informer, discuter avec des amis, partager des informations, exprimer leurs émotions ou se divertir. Un Français sur cinq environ utilise Twitter de manière active ce qui fait de cette plateforme le cinquième réseau social le plus important dans l'hexagone derrière Facebook, Instagram, Snapchat et TikTok¹. Nous nous proposons ici de construire un ensemble d'indicateurs à haute fréquence du moral des Français, à partir des messages postés sur Twitter par un échantillon de plusieurs millions d'utilisateurs. En tant que source, Twitter présente un intérêt particulier : celui de fournir une expression directe et volontaire de l'état d'esprit des internautes. À la différence des enquêtes d'opinion qui interrogent des personnes sur des sujets qu'elles n'ont pas choisis, ou des études qui cherchent à interpréter le comportement des individus, les messages postés sur Twitter sont des expressions actives et délibérées des idées et des sentiments que les personnes ont envie de rendre publiques. Certes, de ce fait, on ne connaît que ce que les intéressés ont souhaité communiquer. Il faut donc considérer ce corpus de tweets en tant que tel et l'interpréter comme un ensemble de manifestations intentionnelles de la population, en réaction aux aléas de la vie économique, sociale et politique.

L'intérêt des chercheurs pour les traces digitales laissées par les internautes n'est pas récent. Depuis une dizaine d'années, des chercheurs en sciences sociales, en économie, en finance, en psychologie et en informatique les exploitent pour construire des indicateurs à haute-fréquence et faire du *nowcasting*. Chaque jour, en naviguant sur Internet, en recherchant

1. DIGIMIND, *Tendances 2021. Les chiffres essentiels pour comprendre les réseaux sociaux*. <https://www.digimind.com/fr/ressources/datamind-2021-chiffres-essentiels-reseaux-sociaux>

des informations sur Google, en lisant des pages sur Wikipédia ou en publiant des messages sur des blogs, forums de discussions ou bien encore sur les réseaux sociaux, les individus laissent des traces de leurs passages qui pourraient permettre de saisir leurs opinions et pensées du moment². Ces nouvelles données permettent de dépasser certaines limites des indicateurs construits à partir de sondages : les données peuvent être disponibles à haute-fréquence (journalière, voire infra-journalière), le coût de collecte est limité, voire nul, et cette approche consistant à « écouter » plutôt qu'à « demander » permet de réduire le biais de non-réponse souvent très important dans les données de sondage³⁴. Bien entendu, la construction de ce nouveau type d'indicateurs soulève un grand nombre de questions en ce qui concerne la non-représentativité susmentionnée de l'échantillon des utilisateurs en ligne, la possibilité d'estimer le bien-être d'un individu à partir de ses traces digitales (biais de désirabilité sociale), ou bien encore les risques de manipulation de ces indicateurs si ces derniers venaient à être utilisés plus directement dans le cadre du choix ou de l'évaluation de politiques publiques. Gardant ces limites à l'esprit, nous tirons parti du fait que la plateforme Twitter est ouverte aux chercheurs et rend ses données publiques, afin de mesurer le « pouls de

2. Ces sources de données sont principalement utilisées pour mesurer le bien-être subjectif, mais peuvent l'être aussi pour mesurer le bien-être objectif. Voir V. Voukelatou, L. Gabrielli, I. Miliou, S. Cresci, R. Sharma, M. Tesconi et L. Pappalardo, *Measuring objective and subjective well-being : dimensions and data sources*, 2021.

3. S. Iacus, G. Porro, S. Salini et E. Siletti, « Controlling for selection bias in social media indicators through official statistics : A proposal », 2020.

4. Il existe aussi d'autres biais : le risque de réponse biaisée (lié notamment au choix de l'ordre des questions) ou encore le problème lié au « Hawthorne Effect » (voir Angus Deaton, « The financial crisis and the well-being of Americans », 2012) : être observé quand on nous demande notre avis sur notre bonheur personnel tend à changer notre comportement. Voir, par exemple, D. Kahneman et A. B. Krueger, « Developments in the measurement of subjective well-being », *Journal of Economic Perspectives*, 20 (1), 2006.

la nation » et le moral des Français. Les métriques d'émotions, ainsi élaborées à partir des messages postés sur le service de *microblogging* Twitter, n'ont pas vocation à remplacer les indicateurs de bien-être calculés à partir des enquêtes, mais plutôt à compléter ces dernières de manière à enrichir notre compréhension des évolutions du bien-être en France.

Ce travail s'inscrit dans le cadre de l'Observatoire du bien-être du Cepremap. Il s'agit, à notre connaissance, du premier projet de cette ampleur sur la base de données françaises. L'ensemble des indicateurs et des informations présentés dans cet opusculé sont disponibles gratuitement en téléchargement et visibles dans le *Tableau de bord du bien-être en France* de l'Observatoire du bien-être.

Avant de présenter les séries élaborées, nous commencerons par décrire la méthodologie utilisée pour constituer notre base de données originale en évitant les robots et les spams, pour qualifier la teneur des sentiments et des émotions des messages et pour reconnaître les préférences politiques, la tranche d'âge et le genre de chaque utilisateur afin de caractériser l'évolution du « moral » de différents sous-groupes de la population.

Notons que nous nous alignons sur la terminologie en vigueur dans ce champ de recherche et employons le terme « sentiment » pour désigner la *valence nette* des messages, c'est-à-dire la teneur plutôt positive ou négative de leurs termes.

Notre score de « sentiment » fondé sur l'analyse textuelle des messages et sur les emojis suit la chronique des chocs positifs et négatifs qui affectent le pays, plongeant ainsi au moment de la crise des Gilets jaunes, puis des confinements et, plus récemment, lors du déclenchement de la guerre en Ukraine, le 24 février 2022. Il présente également de grands écarts autour de la tendance, notamment les jours des attentats de *Charlie Hebdo*, de l'Hyper Casher ou du Bataclan. À l'inverse, les moments de sentiment positif sont moins visibles, à l'exception régulière des jours de fête tels que les veilles et jours de Noël et du Nouvel An.

Par rapport au niveau déprimé de 2020, l'année 2021 a vu une augmentation du sentiment de joie et une diminution du sentiment de tristesse, même si aucun de ces deux indicateurs n'a retrouvé son niveau antérieur à la crise des Gilets jaunes. En revanche, la peur et la colère restent à des niveaux élevés, avec un pic du sentiment de colère, en juillet 2021, à l'annonce de la mise en place du passe sanitaire.

Le score de sentiment connaît aussi des fluctuations régulières au cours de la semaine, et même de la journée. Les utilisateurs de Twitter partent d'un point bas le lundi, puis leur sentiment s'améliore au cours de la semaine pour atteindre un sommet le vendredi et le samedi. À l'échelle de la journée, le sentiment de tristesse, comme celui de peur, atteint un maximum aux petites heures du matin, entre 3 heures et 5 heures, et les expressions de colère et de dégoût connaissent un pic matinal vers 7 heures. Mais la fin d'après-midi et la soirée affichent des émotions plus positives. Les horaires de travail sont plus neutres émotionnellement, tandis que l'on s'exprime l'après-midi et en soirée sur des sujets plus personnels et surtout, plus réjouissants.

Si, dans l'ensemble, on ne distingue pas de différence nette entre les scores émotionnels des hommes et des femmes, on relève cependant quelques différences. Le point le plus haut du score des hommes sur la période correspond à la Coupe du monde masculine de football en 2018 (la victoire de l'équipe de France), sans que l'on puisse distinguer un effet sur le score des femmes. Quant à ces dernières, leur score émotionnel s'est plus dégradé que celui des hommes au moment des confinements imposés par la pandémie.

Enfin, grâce à la représentation du score de sentiment selon les orientations politiques, on observe une séparation de plus en plus nette entre les abonnés aux comptes des partis extrêmes d'un côté et aux partis du centre (EELV, LR, LREM, PS) de l'autre. Au tout début de la période étudiée, en 2013, les abonnés à l'extrême-gauche exprimaient des sentiments plus négatifs. Puis, les sentiments exprimés par les abonnés à des

figures de l'extrême droite se sont dégradés au cours de l'année 2014, venant occuper une position intermédiaire entre l'extrême gauche et les partis traditionnels. L'année 2018 a marqué un tournant : les abonnés à des comptes d'extrême droite ont exprimé au cours de l'année des sentiments de plus en plus négatifs, rejoignant le niveau des abonnés à l'extrême gauche, dont le sentiment était par ailleurs également en baisse. Les abonnés aux comptes de la gauche traditionnelle ont connu, eux aussi, une baisse de leurs sentiments exprimés, mais à partir d'un niveau beaucoup plus élevé. Le fort écart entre les sentiments des extrêmes et ceux des abonnés aux partis traditionnels reste élevé et constant depuis 2019.

Si l'on classe les internautes en fonction des médias qu'ils suivent sur Twitter, on aboutit à un constat similaire. Les personnes qui suivent des médias proches des extrêmes du spectre politique, ici *L'Humanité* et *Valeurs actuelles*, expriment les émotions les plus négatives. Au-delà de cette forte polarisation, les personnes suivant les comptes de médias plutôt à gauche (*Libération*, *Le Monde*, *L'Obs*) ont tendance à exprimer des émotions plus négatives que celles suivant des médias plutôt à droite (*Le Figaro*, *Challenges*).

Nous relevons enfin un effet de génération au sens où les abonnés, généralement jeunes, des deux médias web que sont *Hugo Décrypte* et *Thinker View* expriment des émotions très négatives, au même niveau que les personnes suivant *L'Humanité* ou *Valeurs actuelles*.

Notre indice met donc en évidence une polarisation multiple de l'espace des émotions : d'une part, le contraste politique bien identifié entre le centre et les extrêmes, ces derniers étant caractérisés par un bien-être plus faible et des émotions plus négatives et, d'autre part, un contraste générationnel, avec des médias jeunes associés à une plus forte expression de sentiments négatifs. Ce mal-être a sans doute trouvé une traduction dans les urnes au premier tour de l'élection présidentielle de 2022.

1. Comment mesurer le moral des Français en ligne ?

Poussés par les progrès des méthodes d'analyse textuelle, par l'importance grandissante des traces digitales laissées par les individus en ligne, et par l'ouverture (partielle) des données de certaines grandes plateformes privées, de nombreux chercheurs se sont intéressés à la construction de nouveaux indicateurs de bien-être, d'état émotionnel ou de préoccupation des internautes. Les principales sources de données analysées comportent les requêtes faites par les individus sur Google, les statuts publiés sur Facebook et les messages envoyés sur Twitter. Nous nous attacherons ici à analyser ces messages.

La méthodologie utilisée pour dériver des indicateurs quantitatifs de sentiment à partir des données Twitter peut se décomposer de la manière suivante : (1) constituer une base de documents (tweets) avec des méthodes de *web scraping* ou *via* l'utilisation d'interface de programmation (API) ; (2) identifier le sentiment et les émotions de chaque tweet à partir de méthodes d'analyse textuelle ; (3) construire des indicateurs de sentiment par période/localisation en agrégeant le sentiment de différents tweets ; (4) utiliser ces nouveaux indicateurs pour prévoir, comprendre ou expliquer d'autres variables du monde réel (résultats des élections, évolution des marchés financiers, bien-être des individus, etc.)⁵.

5. Les données Twitter ont été utilisées – avec plus ou moins de succès – dans de nombreux autres domaines tels que la prévision de l'évolution des marchés financiers (P. Azar et A. Lo, « The wisdom of Twitter crowds : Predicting stock market reactions to FOMC Meetings *via* Twitter feeds », *Journal of Portfolio Management*, 42, 2016), la prévision des résultats aux élections présidentielles (A. Tumasjan et al., « Predicting elections with Twitter : What 140 characters reveal about political sentiment », in *Proceedings of the International AAAI Conference on Web and Social Media*, 2010), la propagation d'épidémies (A. Culotta, « Detecting influenza outbreaks by analyzing Twitter messages », *arXiv*, 27 juillet 2010) ou bien encore la mesure de l'incertitude économique (D. Altig et al., « Economic uncertainty before and during the Covid-19 Pandemic », NBER, Working Paper 27418, juin 2020).

CONSTRUIRE UNE BASE DE DONNÉES DE TWEETS

La première étape – avant l'analyse de sentiment ou la construction de nouveaux indicateurs – consiste à construire une base de données de messages. Depuis sa création en 2006 et jusqu'en mi-2020, l'accès aux données Twitter pour les chercheurs était très limité. Il était possible de récupérer, pour un mot-clé donné, un historique des tweets des derniers jours, mais il était impossible de constituer, *via* l'API officielle, des jeux de données sur longue période⁶. La constitution de bases de données historiques n'était possible que par une extraction en temps réel des messages, ce qui limitait donc la possibilité de constituer des séries sur longue période et rendait complexe la répliquabilité des études réalisées sur le sujet⁷. La pandémie de la Covid-19 a entraîné une première modification des conditions d'accès aux données de la part de Twitter avec l'ouverture, en avril 2020, d'un jeu de données de tweets contenant des mots-clés en lien avec la pandémie qui permettent de faciliter la recherche académique⁸. Ce changement s'inscrit dans un cadre plus général d'ouverture des données de la part de grandes entreprises américaines durant la période Covid, par exemple les données de mobilité ouvertes temporairement par Google et Apple. Par la suite, en janvier 2021, Twitter a ouvert un accès gratuit – pour les chercheurs – à l'historique

6. Une API (*Application Programming Interface*) est une interface de programmation d'application. Il s'agit d'un ensemble de règles permettant à des programmes informatiques (ici ceux utilisés par les scientifiques) d'envoyer des demandes normées à un autre programme informatique (ici, les systèmes d'information de Twitter). Ces règles se différencient du *web scraping* par lequel un programme récolte des informations conçues pour être consultées par un humain.

7. Il était possible aussi d'acheter ces bases de données – mais cela était excessivement coûteux – ou bien de collecter les données en passant par des outils non officiels mais ne respectant pas les conditions d'utilisation de Twitter et *via* lesquels la qualité des données récoltées était très imparfaite.

8. Voir A. Tornes, « Enabling study of the public conversation in a time of crisis », *Twitter Development Platform Blog*, avril 2020.

complet des conversations publiques, un accès autrefois réservé aux clients « premium » ou aux entreprises. Ce changement majeur dans la politique d'accessibilité des données par Twitter offre de nombreuses nouvelles possibilités. Dans le cadre de l'accès académique, les données sont accessibles *via* une interface de programmation (API) qui permet d'extraire jusqu'à 10 millions de tweets par mois. De multiples filtres sont disponibles pour collecter des tweets à partir d'un ou plusieurs mots-clés définis par l'utilisateur. Il est également possible de filtrer en indiquant une localisation, une langue ou un intervalle de temps.

Cette limite de 10 millions de messages par mois est importante, mais ne permet pas d'extraire l'ensemble des messages publiés sur Twitter (environ 500 millions de messages par jour). Il est cependant possible de tirer de manière aléatoire un échantillon de tweets en spécifiant des intervalles de temps de quelques secondes chaque jour et en indiquant une langue et/ou en utilisant un mot-clé commun (*stop words*) tel que « le » ou « la » dans la requête. Dans le cadre de l'analyse du bien-être, et comme proposé par C. Yang et P. Srinivasan⁹, il est possible d'utiliser les mots-clés « *I* », « *I'm* », « *my* » et « *mine* » (ou la traduction de ces mots-clés dans une autre langue) ainsi que des intervalles de temps aléatoires afin de constituer un échantillon de tweets.

Nous avons utilisé l'API Twitter pour extraire un échantillon de tweets publiés entre 2010 et 2022 en langue française. Plus précisément, nous avons extrait l'ensemble des tweets contenant le mot-clé « je » publiés durant les soixante premières secondes de chaque heure. Cette méthodologie nous a permis de construire un échantillon aléatoire de plus de 23 millions de messages sur longue période – tout en respectant les limites d'utilisation de l'accès académique à l'API Twitter, soit 10 millions de messages par mois. Nous avons choisi le mot-clé « je » afin de sélectionner principalement des messages dans lesquels les utilisateurs

9. C. Yang et P. Srinivasan, « Life satisfaction and the pursuit of happiness on Twitter », 2016.

expriment leurs actions, opinions et pensées personnelles, comme recommandé par C. Yang et P. Srinivasan. Tout du moins, ce que les utilisateurs de notre échantillon ont envie d'exprimer sur les réseaux sociaux : ce qui n'est pas nécessairement équivalent à ce qu'ils pensent au fond d'eux-mêmes, mais qui – en moyenne sur un ensemble d'utilisateurs – doit tout de même être fortement corrélé avec leurs opinions, émotions, sentiment et envies. La figure 1 montre l'évolution du nombre de tweets durant l'ensemble de la période ainsi que la répartition du nombre moyen de tweets en fonction de l'heure de la journée. Notre base de données est donc un échantillon contenant environ 1/60 de l'ensemble des tweets contenant le mot-clé « je ».

Comme nous pouvons le voir sur la figure 1, le nombre de tweets a largement augmenté de 2010 à mi-2013, avant de connaître une baisse marquée de mi-2013 à 2017, puis de fortement réaugmenter au moment du confinement. Cette augmentation durant le confinement est liée au désir de nombreuses personnes, confinées et isolées, de garder contact avec leurs proches et de « passer le temps » : Twitter a par exemple connu son record d'utilisation durant la pandémie de la Covid-19, selon un article du *Washington Post*¹⁰. L'analyse de la fréquence des tweets en fonction de l'heure de la journée montre que l'utilisation la plus importante de Twitter – pour ce qui concerne les tweets en français contenant le mot-clé « je » – peut être observée entre 18 heures et 1 heure du matin. La nuit – entre 2 heures et 6 heures – le volume de tweets est plus faible.

Les 23 millions de messages de notre échantillon ont été envoyés par un total de plus de 3 millions d'utilisateurs : soit une moyenne d'environ sept tweets par utilisateur. Cependant, une première analyse montre qu'une cinquantaine d'utilisateurs de notre échantillon ont envoyé plusieurs milliers de messages : un nombre très important étant donné que

10. « Twitter sees record number of users during pandemic, but advertising sales slow », *The Washington Post*, 30 avril 2020.

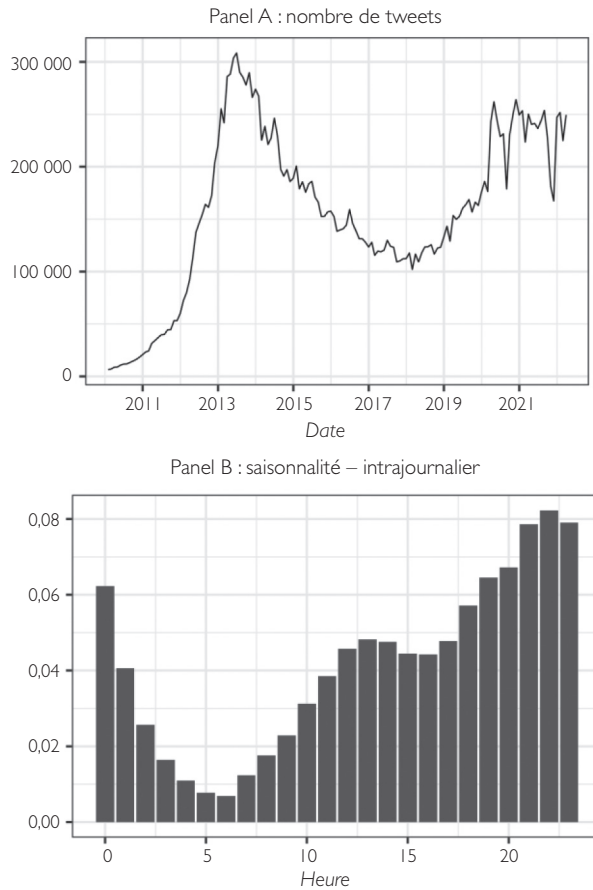


Figure 1 – Nombre de tweets (échantillon de tweets en français).

notre échantillon ne contient qu'un soixantième de l'ensemble des tweets et uniquement des tweets contenant le mot-clé « je ». Par exemple, l'utilisateur ayant envoyé le plus de tweets dans notre échantillon – et qui publie

ses messages sous le pseudo « Afalistolec » – est en réalité un programme informatique (un *bot*, ou robot) qui publie chaque heure un message énigmatique dans le cadre d'une performance artistique¹¹. Nous utilisons alors pour ce compte – comme pour tous les comptes ayant envoyé un nombre important de messages – l'outil Botometer de C. Davis *et al.* (2016) afin de supprimer de manière automatique les *bots*. Dans notre exemple, l'utilisateur Afalistolec obtient un score Botometer de 4,5/5 ce qui indique qu'il s'agit bien d'un robot. Nous calculons – en plus du Botometer – un score de similarité de tweets pour chaque utilisateur dans le but de supprimer les spams qui ne sont pas toujours parfaitement détectés par Botometer, et un filtre en prenant en compte la présence de liens dans les tweets¹² ayant montré que les robots utilisent beaucoup plus de liens dans les messages que les humains¹³. L'application de ces trois filtres permet la suppression d'environ 1,5 million de messages envoyés par un total de 440 robots identifiés, soit environ 5 % de l'ensemble des tweets.

MESURER LES SENTIMENTS ET LES ÉMOTIONS

La seconde étape qui conduit à l'analyse de sentiment consiste à convertir une variable qualitative, un tweet par exemple, en variable quantitative. La classification la plus commune consiste à identifier pour chaque texte

11. L'auteur de cette performance est Rico de Halvarez. Le programme en question publie des extraits de textes ; voir <https://actonart.fr/rico-da-halvarez/>

12. S. Stieglitz, F. Brachten, D. Berthélé, M. Schlaus, C. Venetopoulou et D. Veutgen, « Do social bots (still) act different to humans ? Comparing metrics of social bots with those of humans », 2017.

13. Ces filtres permettent de supprimer, par exemple, les messages d'un utilisateur ayant envoyé 5 000 fois un tweet contenant le mot-clé #galaxys6 afin de participer à un concours pour gagner un téléphone. Si cet utilisateur n'est pas supprimé, du fait de la présence du mot-clé « gagner » dans les tweets relatifs à ce concours, les indicateurs de sentiment agrégé seront fortement biaisés durant les mois de mars, avril et mai 2015.

une polarité -1/+1 (positif/négatif) ou -1/0/+1 (positif/neutre/négatif) et ce, à partir d'un dictionnaire de mots pré-classifiés comme étant positif ou négatif. De nombreux dictionnaires de mots positifs/négatifs sont disponibles – en anglais en très grande majorité – et permettent de calculer rapidement un score de sentiment pour un document donné¹⁴.

Dans le cadre de cet opuscule, notre objectif étant de quantifier le moral des Français à partir de tweets en français, nous utilisons les deux méthodes suivantes. Tout d'abord, nous employons la méthode VADER¹⁵, spécifique à l'analyse de sentiment des messages publiés sur les réseaux sociaux et dont une implémentation pour l'analyse de tweets en français existe depuis mai 2020¹⁶. La méthode VADER présente l'avantage de prendre en compte les emojis, les négations, la casse et la ponctuation – en plus des mots contenus dans les tweets – afin de calculer un score de sentiment pour chaque message (voir le tableau I pour des exemples de texte et le score calculé par VADER dans chacun des cas). Pour chaque tweet, nous utilisons le score de sentiment de VADER afin de calculer pour chaque message un score entre -1 et +1 (de très négatif à très positif).

14. Étant donné la plus grande disponibilité de packages et d'outils d'analyse textuelle en anglais, il est aussi possible de traduire les tweets vers l'anglais – *via* DeepL ou Google Translate – afin d'appliquer les méthodes d'analyse développées pour les textes en anglais. Cette méthode est cependant coûteuse – la traduction automatique d'un nombre important de tweets est payante – et convient principalement à l'analyse de tweets dans des langues pour lesquelles il n'existe aucun outil, ce qui n'est pas le cas pour le français.

15. C. Hutto et E. Gilbert, « VADER : A parsimonious rule-based model for sentiment analysis of social media text », 2014.

16. <https://pypi.org/project/vaderSentiment-fr/>

Tableau 1 – Exemples de tweets classifiés avec l’approche VADER

Tweet	Score de sentiment
Je suis heureux	0.5719
Je suis trop heureux	0.6115
Je suis TROP heureux	0.7200
Je suis TROP heureux 😊	0.9070
Je suis TROP heureux 😊 !!!	0.9219
Je ne suis pas heureux	– 0.4585
Je suis triste 😞	– 0.4767
Je suis triste	– 0.6705
J’en ai marre	– 0.4215
J’en ai jamais marre	0.3252

Pour la mesure des émotions, la méthode FEEL – *French Expanded Emotions Lexicon* – permet de classifier chaque message selon les six émotions : la joie, la colère, la peur, la tristesse, le dégoût et la surprise¹⁷. Cependant, ce dictionnaire ayant été développé pour analyser des textes traditionnels, il ne permet pas de capturer les spécificités des messages publiés sur les réseaux sociaux, notamment parce qu’il ne contient pas d’emojis. C’est pourquoi nous privilégions dans cet opuscule une approche fondée sur l’analyse des emojis, en utilisant EmoTagI200 – une liste de 150 emojis classifiés en fonction de 8 émotions¹⁸. Nous détaillons ces choix et les différentes mesures dans l’Annexe I. Sur l’ensemble de notre échantillon, 13 % environ des messages envoyés contiennent

17. A. Abdaoui, J. Azé, S. Bringay et P. Poncelet, « FEEL : A French expanded emotion lexicon », *Language Resources & Evaluation*, 51 (1), 2017.

18. A. Shueb et G. De Melo, « EmoTagI200 : Understanding the association between emojis and emotions », 2020.

au moins un émoji : un nombre équivalent à celui calculé par Emojipedia à partir de l'analyse de plus 6 milliards de tweets envoyés partout dans le monde¹⁹. Pour construire nos indicateurs, et afin de contrôler pour la forte variation du nombre de tweets contenant des émojis dans le temps et pour la variation de l'émotivité, nous divisons pour chaque période le score d'émotions calculé à partir de EmoTag1200 par l'émotivité de l'ensemble des tweets (la somme des scores associés à toutes les émotions).

19. <https://blog.emojipedia.org/emoji-use-at-all-time-high/>

2. Le moral des Français au prisme de Twitter

LE MORAL GLOBAL DES FRANÇAIS SUR TWITTER : DU MOIS À LA JOURNÉE

La figure 2 présente l'évolution de notre indicateur sur l'ensemble de la période – en considérant VADER pour l'analyse de sentiment et en agrégeant par période en prenant la moyenne du score de sentiment de l'ensemble des tweets – après suppression des messages envoyés par des robots et ajustement pour les variations saisonnières pour corriger du score de sentiment anormalement élevé chaque année au moment de Noël, du Nouvel An, de la Saint-Valentin ainsi que durant les vacances. Notre indicateur est centré (moyenne 0) et standardisé afin de faciliter la lecture.

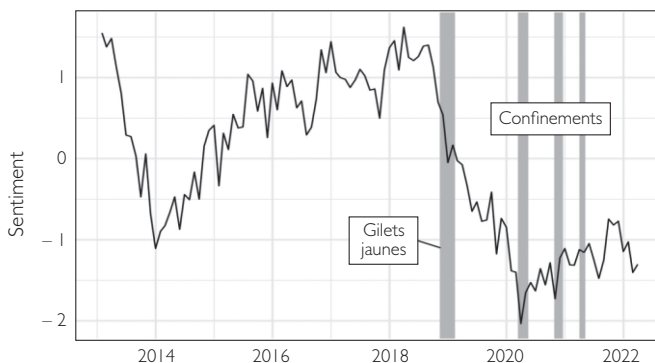


Figure 2 – Indicateur Twitter du moral des Français (fréquence mensuelle).

Note de lecture : les zones grises sur le graphique correspondent respectivement au début des Gilets jaunes (novembre 2018), puis aux trois confinements (mars 2020, avril 2021, novembre 2021).

Durant la période étudiée, nous observons une chute rapide du score de sentiment entre 2013 et 2014, suivie d'une progression irrégulière de 2014 à 2019. Cette augmentation révèle qu'à cette aune, des

événements tragiques comme les attentats de janvier (*Charlie Hebdo*) et de novembre 2015 n'ont eu qu'un impact ponctuel sur les sentiments exprimés (voir aussi la figure 4 ci-après). À contrario, le graphique fait apparaître assez nettement la dégradation des sentiments exprimés à partir de l'automne 2018, prélude à la crise des Gilets jaunes. Si les occupations de ronds-points et les manifestations ont progressivement diminué en intensité au cours de l'année 2019, l'indicateur suggère que les sentiments négatifs qui ont propulsé une partie de cette crise n'ont pas disparu pour autant. On voit certes une amélioration fin 2019, mais on entre ensuite en plein dans la pandémie et ses conséquences. En particulier, les entrées dans le premier et le deuxième confinement ont marqué une nouvelle dégradation des sentiments exprimés. Ce n'est qu'à la sortie du deuxième confinement que l'indicateur se stabilise quelque peu à son niveau, déjà déprimé, d'avant la pandémie. À l'échelle du mois, notre indicateur met donc en relief les principales tendances de l'état émotionnel des Français, qui évolue à cet horizon en fonction des grands chocs qui traversent la société.

Un intérêt important des données issues des réseaux sociaux tient à leur fréquence : on peut regarder ce qui se passe à l'échelle du mois, mais aussi de la semaine ou de la journée. L'évolution du score de sentiment au jour le jour est présentée dans la figure 3. Le sentiment est, on le voit, volatil : autour de la tendance, les écarts d'un jour à l'autre sont importants. La plupart des jours, ils restent dans une bande de variation autour de la tendance. Plusieurs dates se détachent toutefois nettement par un sentiment beaucoup plus négatif. Ces jours sont le 7 janvier 2015 (attentats de *Charlie Hebdo* et de l'Hyper Casher), le 14 novembre 2015 (lendemain des attentats du Bataclan), le 15 juillet 2016 (attentat de Nice), le 26 janvier 2020 (mort de Kobe Bryant), les 15 et 16 mars 2020 (annonce du premier confinement) et le 24 février 2022 (guerre en Ukraine).

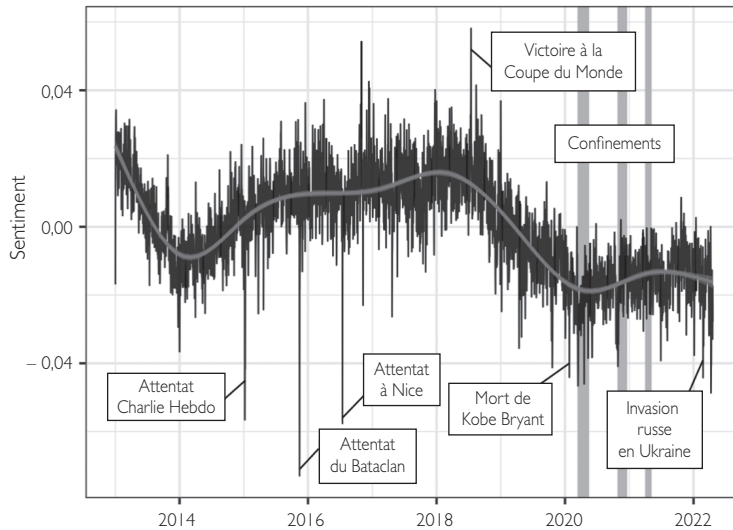


Figure 3 – Indicateur Twitter du moral des Français (fréquence journalière).

À l’opposé, les moments de sentiment positif sont moins directement visibles. En pratique, et c’est un constat qui avait déjà été fait sur les statuts Facebook, les jours les plus heureux sont les jours de fêtes : veille/ jour de Noël ou veille/jour du Nouvel An²⁰. Une fois corrigé de ces variations saisonnières, le jour ayant connu le score de sentiment le plus élevé sur l’ensemble de notre période d’étude est le 15 juillet 2018 : jour de la victoire de l’équipe de France de football en Coupe du monde.

20. Thanksgiving Day est aux États-Unis le jour où le sentiment est le plus élevé. En France, il s’agit du Nouvel An. Un biais possible concernant ces jours est la très forte présence de messages tels que « bonne année » ou « joyeux Noël » qui tendent à être classifiés positivement dans les méthodes d’analyse de sentiment.

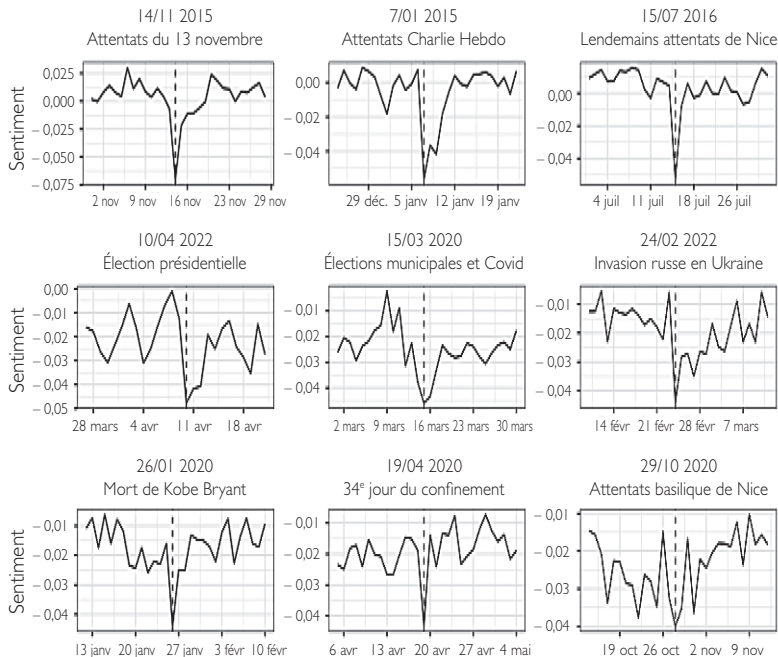


Figure 4 – Évolution du score de sentiment autour des chocs négatifs majeurs.

Un point important à relever est que ces scores extrêmes de sentiment, négatifs comme positifs, sont très ponctuels : on peut les associer directement à une journée, et il semble que l'indicateur retrouve très rapidement son niveau d'avant l'évènement. Les utilisateurs passent-ils aussi rapidement à autre chose, ou est-ce que l'ampleur de l'écart le jour même nous empêche de voir une tendance les jours suivants ? Pour trancher entre ces deux possibilités, nous observons la variation de sentiment autour des 9 principaux chocs négatifs de notre échantillon en analysant le sentiment sur une fenêtre de 15 jours autour de chaque évènement. La

figure 4 présente nos résultats. Elle confirme que ces chutes de score de sentiment sont concentrées sur une journée, ou au plus quelques jours. Avec notre indicateur, nous mesurons donc la réaction *immédiate* à un évènement, la façon dont les utilisateurs actifs de Twitter le reçoivent et l'expriment dans l'instant. La période du premier confinement apparaît toutefois comme un peu particulière. Plutôt qu'un effet immédiat important lors de l'annonce du confinement, nous observons une série plus dense de pics négatifs, et surtout une tendance nettement infléchie à la baisse pendant le confinement. Cela vient confirmer que celui-ci a été une épreuve pour beaucoup, un élément que nous retrouvons dans les données d'enquête sur le bien-être émotionnel.

SENTIMENTS ET ÉMOTIONS

L'indicateur de sentiment agrégé est un indice synthétique, fondé sur une analyse du texte. L'analyse des émotions à partir des émojis offre une grille de lecture complémentaire à l'analyse de sentiment à partir de VADER. La figure 5 présente les indicateurs par émotions – ajustés du nombre total d'émojis pour chaque période. Il apparaît immédiatement que l'indice de sentiment évolue globalement de la même manière que la série qui indique la fréquence des expressions de la joie, et à l'inverse de celles qui expriment la tristesse, la peur ou la colère.

Cette approche en plusieurs sentiments (émotions) permet de constater que par rapport au niveau déprimé de 2020, 2021 a vu une augmentation du sentiment de joie et une diminution du sentiment de tristesse. On se rapproche ainsi fin 2021 des niveaux d'avant la pandémie. Ce n'est qu'une maigre consolation : la joie reste moins fréquente qu'avant la crise des Gilets jaunes et la tristesse, plus fréquente. En revanche, la peur et la colère restent à des niveaux élevés. Le pic du sentiment de colère, en juillet 2021, correspond à l'annonce de la mise en place du passe sanitaire et, avec lui, d'une forte incitation à la vaccination.

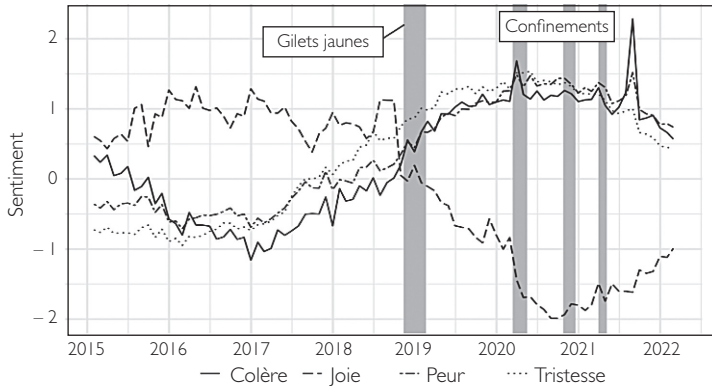


Figure 5 – Évolution des indicateurs d'émotion (fréquence mensuelle).

LES HEURES ET LES JOURS

Il y a des moments de la journée et des jours de la semaine plus propices aux sentiments positifs – un phénomène déjà relevé par P. Dodds *et al.*²¹. À l'échelle de la semaine (figure 6), les utilisateurs de Twitter partent d'un point bas le lundi, puis leur sentiment s'améliore au cours de la semaine. Le vendredi et le samedi sont assez nettement les jours où le sentiment est le plus positif. Il faut garder à l'esprit que les personnes qui utilisent Twitter ne sont pas les mêmes tout du long. On peut s'attendre à ce que les utilisateurs en semaine soient à peu près les mêmes d'un jour à l'autre. La pente du lundi au vendredi peut donc être lue comme reflétant un rapport tendu avec le travail. Le pic du vendredi traduit un sentiment très positif du passage en week-end, cohérent avec ce que l'on connaît par ailleurs des activités et des sorties. L'usage de Twitter le week-end constitue en revanche un filtre sans doute plus fort. Certains

21. P. Dodds, K. H. Harris, C. Bliss et C. Danforth, « Temporal patterns of happiness and information in a global social network : Hedonometrics and Twitter », 2011.

profils – typiquement les parents ou les personnes utilisant Twitter dans le cadre de leur activité professionnelle – sont beaucoup moins présents sur le réseau le week-end, ce qui modifie la composition de l'échantillon. À ce titre, la baisse du dimanche peut traduire simplement le fait que les personnes actives sur Twitter ce jour-là ont moins que les autres des activités en famille ou en groupe, ce qui en fait un public à priori moins propice à l'expression de sentiments positifs.

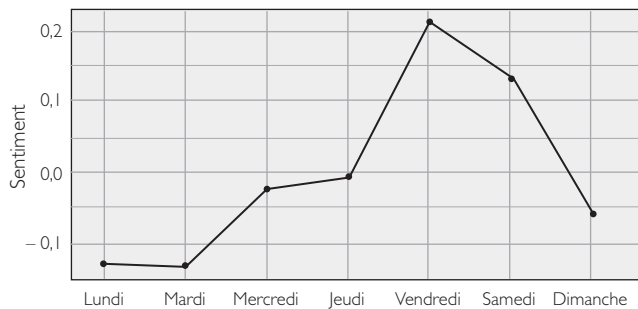


Figure 6 – Score de sentiment moyen en fonction du jour de la semaine.

À l'échelle de la journée, nous observons aussi des variations importantes selon l'heure, probablement liées là aussi à *qui* utilise Twitter à un moment donné – un élément déjà relevé par S. Golder et M. Macy²². Le sentiment de tristesse, comme celui de peur, atteint un maximum aux petites heures du matin, entre 3 heures et 5 heures, et les expressions de joie sont à leur étiage à ce même moment (figure 7). Ce phénomène montre que Twitter sert à certaines personnes de moyen d'expression de leur détresse, particulièrement à des moments où il est difficile de

22. S. Golder et M. Macy, « Diurnal and seasonal mood vary with work, sleep, and day length across diverse cultures », 2011.

parler à d'autres personnes. La colère et le dégoût connaissent de la même manière un pic matinal, décalé vers 7 heures, probablement sous l'influence des personnes qui consultent Twitter au moment de leur réveil.

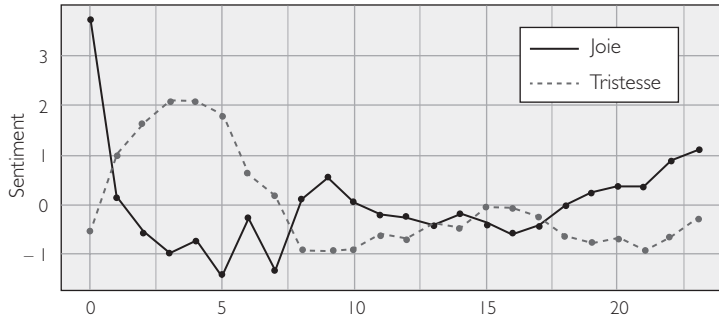


Figure 7 – Variation infra-journalière des émotions.

La fin d'après-midi et la soirée affichent un niveau de sentiment positif plus élevé, et de sentiments négatifs moins élevés. On ne parle pas de la même chose sur Twitter en fonction de l'heure. Les horaires de travail sont plus neutres émotionnellement, tandis que l'on s'exprime l'après-midi et en soirée sur des sujets plus personnels et surtout, plus réjouissants.

GENRE, POLITIQUE ET ÉMOTIONS

Certains profils Twitter affichent clairement des informations les concernant, comme leur genre, leur profession ou leur orientation politique, à l'image des sympathisants de La France insoumise qui utilisent le ϕ comme symbole. Il s'agit toutefois d'une minorité. Pour autant, un survol rapide des échanges permet souvent de se faire une idée des penchants politiques de la personne et souvent de son genre et de sa classe d'âge.

Des travaux, par exemple K. Jaidka *et al.*²³, ont construit des corpus de mots révélateurs de ces éléments, mais ceux-ci n'existent que pour la langue anglaise et sont fréquemment spécifiques au contexte politique et social des États-Unis. Pour autant, il semble fondamental de pouvoir distinguer l'état des sentiments selon ces dimensions afin de comprendre quelles forces contribuent au mouvement d'ensemble de l'état des Français. Nous employons donc une série d'approches alternatives en fonction des dimensions visées : le prénom pour en inférer le genre, les comptes politiques suivis pour les préférences politiques, les médias suivis pour les préférences politiques et l'âge. Ces éléments permettent de dessiner un espace des sentiments structurés à la fois par le genre, par l'orientation politique et par la génération.

LES FEMMES DANS LA PANDÉMIE

Le *Tableau de bord de l'Observatoire du bien-être* du Cepremap²⁴ montre que les femmes connaissent en moyenne une satisfaction dans la vie plus faible que les hommes, une propension à s'être senties heureuses la veille plus forte, mais aussi à se sentir plus déprimées. Ainsi, si le bien-être émotionnel a une dimension genrée, il n'est pas évident à priori de savoir quelle dimension va dominer dans les sentiments exprimés sur Twitter.

Une difficulté : de nombreux utilisateurs et de nombreuses utilisatrices n'indiquent pas leur genre dans leur biographie Twitter. Pour repérer le genre, nous nous limitons aux personnes qui déclarent un prénom, et classifions celui-ci à l'aide d'une liste de 11 627 prénoms identifiés comme féminins et masculins²⁵. Nous sommes ainsi en mesure de calculer sur

23. K. Jaidka, S. Giorgi, S., H. A. Schwartz, M. Kern, L. Ungar, et J. Eichstaedt, « Estimating geographic subjective well-being from Twitter : A comparison of dictionary and data-driven language methods », 2020.

24. <http://www.cepremap.fr/Duree.html>

25. La base <https://www.data.gouv.fr/fr/datasets/liste-de-prenoms/> mobilise les données de Frantext <https://www.frantext.fr/>, corpus de textes essentiellement littéraires.

cet échantillon réduit les indicateurs de sentiment en fonction du caractère masculin ou féminin du prénom que les utilisateurs et utilisatrices de Twitter ont choisi de se donner²⁶.

La figure 8 présente l'évolution de ces deux indicateurs. Dans l'ensemble, les deux séries sont proches et évoluent de manière parallèle. Il n'y a le plus souvent pas de différence massive entre hommes et femmes. Nous relevons toutefois un ensemble de divergences révélatrices. Le point le plus haut du sentiment des hommes sur la période correspond à la Coupe du monde masculine de football en 2018, qui a vu la victoire de l'équipe de France. Cela s'accompagne d'une amélioration assez durable des émotions exprimées par les hommes, sans que l'on puisse distinguer un tel effet auprès du genre féminin.

En outre, les femmes semblent avoir été plus affectées par les deux premiers confinements que les hommes. Leur sentiment déclaré se détériore plus que celui de ces derniers. Malgré une amélioration sensible à la suite du second déconfinement, il reste jusqu'à la fin de la période nettement inférieur à celui des hommes. Cette observation est cohérente avec de nombreuses études (voir par exemple C. Nienhuis et I. Lesser²⁷ et T. Halldorsdottir et al.²⁸, pour l'impact sur les adolescentes) et des sondages réalisés durant la période Covid. Tous documentent une hausse de l'anxiété des femmes durant cette période. Sur un plan très

26. Une limite potentielle de cette méthode réside dans le fait que la propension à afficher un prénom peut être différente entre hommes et femmes. Ces dernières peuvent être plus enclines à dissimuler leur genre afin d'éviter le harcèlement ou les messages à caractère sexuel. L'échantillon des femmes serait alors biaisé d'une manière différente de celui des hommes.

27. C. Nienhuis et I. Lesser, « The impact of Covid-19 on women's physical activity behavior and mental well-being », 2020.

28. T. Halldorsdottir, I. Thorisdottir, C. Meyers, B. Asgeirsdottir, A. Kristjansson, H. Valdimarsdottir et I. Sigfusdottir, « Adolescent well-being amid the Covid-19 pandemic : Are girls struggling more than boys ? », 2021.

pratique, les femmes ont eu plus que les hommes à gérer les aléas de la situation, en particulier lorsqu'il a fallu organiser la garde des enfants, les arrêts et les reprises des activités. Dans le domaine professionnel, celles qui étaient en télétravail avaient beaucoup moins souvent que leurs conjoints un espace dédié pour travailler. Début 2022, l'écart est toujours persistant et, en moyenne, les femmes ont un indice de sentiment inférieur d'environ 15 % à celui des hommes.

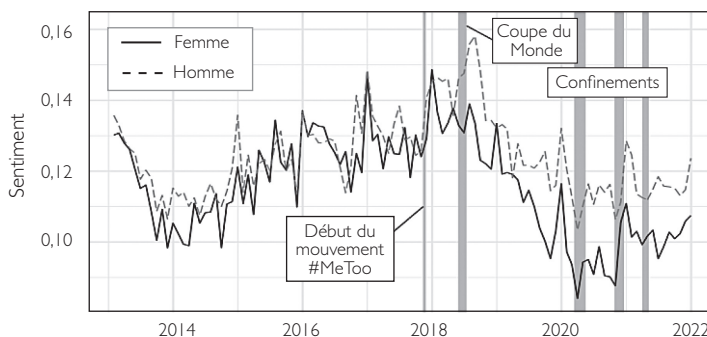


Figure 8 – Indicateur Twitter du moral des Français selon le genre.

LE CENTRE ET LES EXTRÊMES : LA POLARISATION POLITIQUE DES ÉMOTIONS

Sur Twitter, vous devez vous abonner à un compte pour que ses messages arrivent directement dans votre fil d'actualité. Suivre un compte implique donc un choix et une action de la part des utilisateurs, l'expression de la volonté de savoir ce que dit le compte en question. Même si ce n'est pas systématiquement vrai, nombre d'individus vont avoir tendance à suivre les comptes avec lesquels ils sont en accord et partagent des idées. C'est particulièrement le cas pour les comptes de personnalités ou de partis politiques qui servent essentiellement de relais à la communication partisane.

Cette méthode d'analyse des relations a été utilisée par Barbera aux États-Unis afin de dériver les préférences politiques de chaque utilisateur (Démocrate ou Républicain)²⁹. Un algorithme de *machine learning* est entraîné à partir de l'analyse des amis d'une liste de personnalités politiques – dont les préférences politiques sont connues par leur rattachement à un parti – et utilisé ensuite pour classer l'ensemble des autres utilisateurs. J. Golbeck et D. Hansen³⁰ utilisent une méthode semblable pour dériver les préférences politiques à partir des comptes de personnalités politiques suivis par chaque utilisateur et indiquent que cette méthodologie est généralisable à d'autres attributs (âge, revenu) à partir du moment où un jeu d'entraînement d'utilisateurs pour lesquels des caractéristiques sociodémographiques sont connues peut être constitué. P. Ludu propose ainsi d'inférer le genre d'un utilisateur à partir des comptes de célébrités suivies par cet utilisateur³¹.

En ce qui concerne les préférences politiques, nous utilisons une méthode proche de celle de Barbera en analysant les comptes suivis par chaque utilisateur et en estimant les préférences politiques en fonction du nombre et de l'orientation des comptes de personnalités politiques et de partis suivis par cet utilisateur. Nous construisons pour cela une liste d'environ 500 personnalités politiques françaises ayant un compte Twitter, puis nous utilisons cette liste pour classer chaque utilisateur sur une échelle allant de l'extrême gauche à l'extrême droite.

Notre indicateur est construit en analysant *rétrospectivement* les messages des personnes qui suivent – aujourd'hui – un ensemble de comptes politiques³². Nous sommes évidemment conscients que le fait

29. P. Barberá, « Birds of the same feather tweet together : Bayesian ideal point estimation using Twitter data », 2015.

30. J. Golbeck et D. Hansen, « A method for computing political preference among Twitter followers », 2014.

31. P. Ludu, « Inferring gender of a twitter user using celebrities it follows », 2014.

32. Une limite de cette méthode est que nous classifions un utilisateur, et donc l'ensemble de ses messages passés, en fonction de ses abonnements actuels. Cela

de suivre un compte n'implique pas toujours une adhésion aux idées exprimées par ce compte. Il semble cependant plausible que les personnes abonnées au compte officiel d'une personnalité politique soient *majoritairement* des personnes proches des idées de cette dernière. Nous nous concentrons uniquement sur les utilisateurs qui suivent au moins deux politiciens afin de diminuer le bruit lié, par exemple, au fait que plusieurs millions d'utilisateurs suivent le compte d'Emmanuel Macron sans nécessairement soutenir ce dernier. Nous relierons ensuite un utilisateur à une préférence politique seulement si, parmi l'ensemble des comptes politiques qu'il suit, au moins 50 % appartiennent à une même tendance politique. Cette approche permet de représenter la manière dont, au sein de la population française, les émotions exprimées diffèrent en fonction des préférences politiques identifiées à partir des comptes suivis.

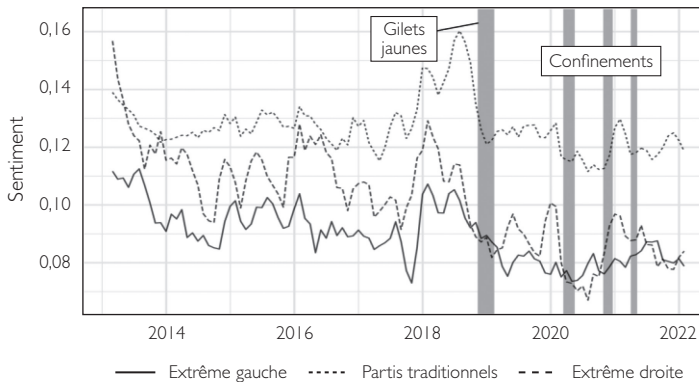


Figure 9 – Moral des Français sur Twitter par préférence politique.

nous conduit certainement à mal classer les personnes dont l'orientation politique a évolué au cours du temps, mais à cette échelle de temps, les orientations des utilisateurs sont certainement assez stables pour que les abonnements actuels reflètent de manière suffisamment solide une orientation politique latente.

La représentation du score de sentiment selon les orientations politiques ainsi inférées se lit sur la figure 9. L'élément le plus significatif est la séparation de plus en plus nette entre les extrêmes d'un côté et les partis du centre (EELV, LR, LREM, PS) de l'autre. Au tout début de nos séries, en 2013, la ligne de partage passe entre une extrême gauche qui exprime des sentiments plus négatifs et les autres partis. Au cours de l'année 2014, les sentiments exprimés par les abonnés à des figures de l'extrême droite se dégradent, venant occuper une position intermédiaire entre l'extrême gauche et les partis traditionnels. L'année 2018 marque un tournant : jusque-là proches des autres partis, les abonnés à des comptes d'extrême droite expriment au cours de l'année des sentiments de plus en plus négatifs, rejoignant le niveau des abonnés à l'extrême gauche, dont le sentiment est par ailleurs aussi en baisse. Les abonnés aux comptes de la gauche traditionnelle connaissent aussi une baisse de leurs sentiments exprimés, mais à partir d'un niveau beaucoup plus élevé. À partir de début 2019, on observe ainsi un écart fort entre les sentiments des extrêmes et ceux des abonnés aux partis traditionnels. En d'autres termes, les sentiments exprimés sont restés significativement différents tout au long des deux années de pandémie, ce qui suggère que les Français n'ont pas vécu de la même manière ces deux ans selon leur orientation politique.

Cette vision recoupe un ensemble de travaux sur la façon dont le bien-être et les émotions structurent le paysage politique français, en particulier l'opposition entre les partis traditionnels et les extrêmes. Y. Algan et *al.* ont ainsi montré que le vote aux extrêmes s'explique largement par une faible satisfaction dans la vie, la différence entre les deux camps se faisant entre un niveau de confiance envers les autres plus élevé à l'extrême gauche et plus faible à l'extrême droite³³.

33. Y. Algan, E. Beasley, D. Cohen et M. Foucault, *Les Origines du populisme, enquête sur un schisme politique et social*, 2019.

Notre indicateur de sentiment indique un changement, entre 2014 et 2018, de l'état des sympathisants d'extrême droite. À la dégradation initiale s'est ajoutée une chute concomitante de l'émergence du mouvement des Gilets jaunes, qui n'a pas été compensée depuis. Ainsi, les proches de l'extrême droite n'ont pas abordé l'élection présidentielle dans le même état émotionnel que lors du scrutin de 2017.

POLARISATION DE L'ESPACE MÉDIATIQUE

Une autre manière d'aborder ce décalage est d'inférer les penchants politiques des utilisateurs de Twitter des médias suivis. Il s'agit d'un vaste chantier d'analyse du paysage médiatique français, dont nous ne présentons ici que quelques éléments exploratoires, que nous développerons dans des travaux ultérieurs. Parmi les 100 comptes les plus suivis de notre échantillon, nous extrayons une liste des médias les plus suivis, que nous classons par type (tableau 2).

Tableau 2 – Médias sur Twitter par type

Quotidiens nationaux	<i>L'Humanité, Libération, Le Monde, L'Équipe, Les Échos, Le Figaro</i>
Quotidiens régionaux	<i>La Voix du Nord, La Provence, Ouest-France, Le Parisien, Sud-Ouest</i>
Hebdomadaires d'information	<i>Le Canard enchaîné, Courrier international, L'Obs, L'Express, Le Point, Le Journal du dimanche, Challenges, Valeurs actuelles</i>
Hebdomadaires culturels et de divertissement	<i>Télérama, Télé loisirs, Elle, Paris Match, Le Figaro Madame</i>
Mensuels	<i>Sciences et Avenir, SoFoot, Vogue, Les Inrockuptibles</i>
Télévision	<i>Quotidien</i>
Web	<i>Hugo Décrypte, Thinker View, Mediapart, Les Jours</i>

Ce regroupement fait apparaître des contrastes marqués selon le type de média (figure 10). Les émotions exprimées par les personnes qui suivent les hebdomadaires culturels et de loisir sont en moyenne nettement plus positives que celles exprimées par les personnes qui suivent

la presse d'information, quotidienne ou hebdomadaire. La presse quotidienne régionale occupe une place intermédiaire.

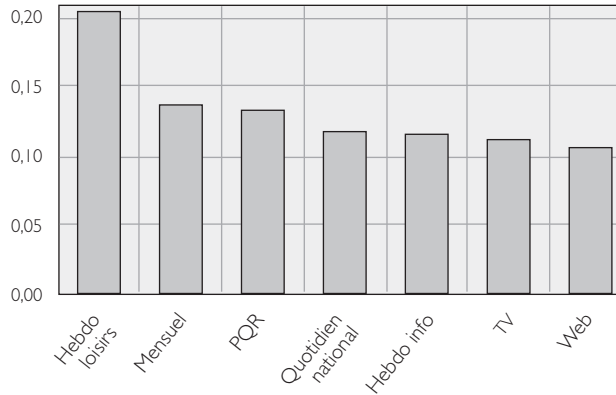


Figure 10 – Indice moyen de sentiment et médias.

Titre par titre, ce regroupement reste apparent, avec d'un côté les hebdomadaires culturels et de loisir, et une grande proximité de la presse quotidienne régionale, à l'exception notable du *Parisien*, qui vient s'intercaler dans le groupe des titres d'information (figure 11).

Par construction, la presse d'information présente une coloration politique plus importante que la presse de loisir. Nous y retrouvons ainsi le constat posé plus haut sur la proximité politique. Les personnes qui suivent des médias colorés aux extrêmes du spectre politique, ici *L'Humanité* et *Valeurs actuelles*, expriment des émotions plus négatives que celles qui suivent des médias plus au centre. Au-delà de cette forte polarisation, nous relevons ici que les personnes suivant les comptes de médias plutôt à gauche (*Libération*, *Le Monde*, *L'Obs*) ont tendance à exprimer des émotions plus négatives que celles suivant des médias plutôt à droite (*Le Figaro*, *Challenges*).

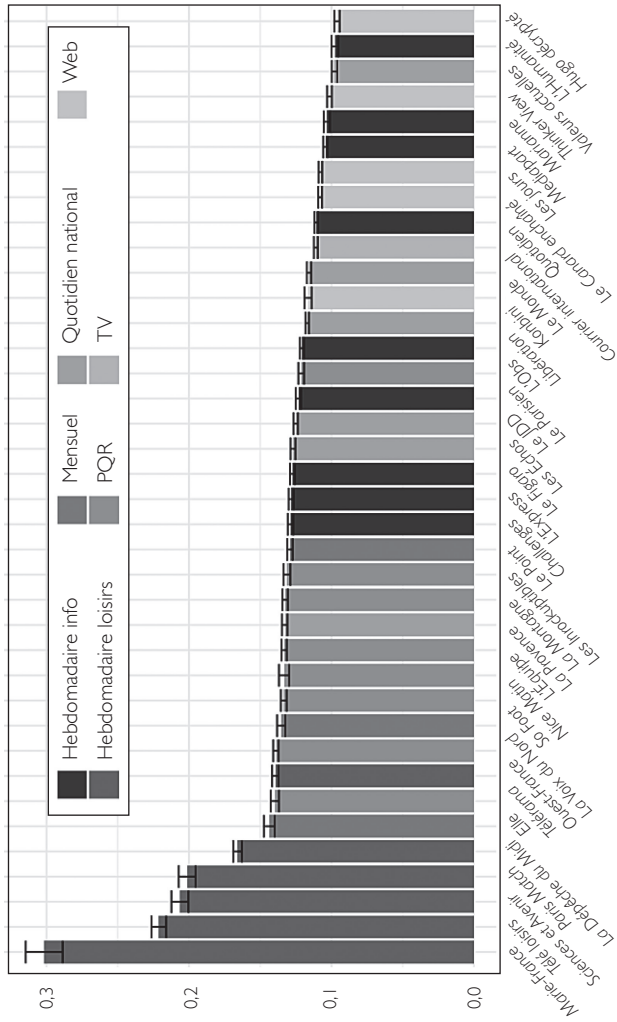


Figure 11 – Indice moyen d'émotions et médias.

Nous remarquons en particulier que les deux médias web ayant une audience plutôt jeune (*Hugo Décrypte* et *Thinker View*) sont parmi les quatre titres ayant les *followers* exprimant les émotions les plus négatives. En moyenne, ils se positionnent au même niveau que les personnes suivant *L'Humanité* ou *Valeurs actuelles*. Il ne s'agit pas ici à notre sens d'un effet de polarisation politique : si *Thinker View* a une orientation qu'on peut qualifier volontiers d'antisystème, *Hugo Décrypte* se situe à peu près au centre du spectre politique en termes de ligne éditoriale. Nous pensons donc que nous relevons là un effet de génération, avec des jeunes qui sont en moyenne plus anxieux et en colère que leurs aînés. Ces deux émotions peuvent avoir des racines et des formes d'expression différentes. *Hugo Décrypte* fait la part belle aux questions d'environnement qui sont particulièrement source d'anxiété chez les jeunes³⁴, tandis que *Thinker View* invite plus volontiers des figures contestataires.

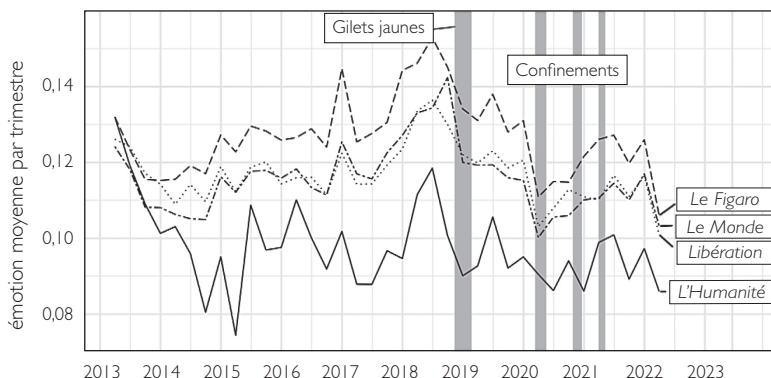


Figure 12 – Presse quotidienne nationale et émotions.

34. Voir par exemple Pascale Krémer et Audrey Garric, « Éco-anxiété, dépression verte ou "solastalgie" : les Français gagnés par l'angoisse climatique », *Le Monde*, 21 juin 2019.

À l'échelle de la presse quotidienne nationale (figure 12), nous voyons que l'écart entre les personnes suivant le compte de *L'Humanité* et les autres apparaît dès 2014. Les évolutions sont ensuite relativement parallèles. Nous retrouvons le pic positif de la Coupe du monde de football masculin en 2018, puis la chute concomitante de l'émergence du mouvement des Gilets jaunes. C'est à ce moment-là, fin 2018 et début 2019, que semble se mettre en place le décalage entre *Le Monde* et *Libération* d'une part, *Les Échos* et *Le Figaro*, d'autre part. Ces deux derniers titres connaissent avec la fin de la phase la plus aiguë du mouvement un regain d'émotions positives, alors que les deux premiers restent à un niveau plus faible. Le premier confinement correspond à une chute brutale pour la quasi-totalité des abonnés, puis une lente reprise au fur et à mesure de la sortie progressive de l'épidémie et des mesures de confinement³⁵. Cette réaction à l'épidémie est d'ailleurs plus marquée chez les personnes suivant des comptes de médias que dans l'ensemble des utilisateurs de Twitter. En d'autres termes, l'exposition – volontaire, puisque ces personnes se sont abonnées aux comptes en question – aux informations engendre une réponse émotionnelle plus orientée que celle du reste de la population.

Du côté des hebdomadaires d'information, nous observons les mêmes dynamiques d'ensemble. Comme pour *L'Humanité*, le décrochage entre la majorité des titres et *Valeurs actuelles* s'opère au tournant de 2014. Les personnes qui suivent *Le Canard enchaîné* expriment également des émotions plus négatives que la moyenne – l'aspect satirique du journal ne semblant pas compenser l'impact négatif des révélations qui font la ligne éditoriale de ce titre (figure 13).

35. Nous ne représentons pas ici la figure correspondante pour la presse quotidienne régionale dans la mesure où les évolutions sont similaires à celles des principaux titres de la presse nationale, et que nous n'observons pas de contrastes marqués entre les différents titres régionaux.

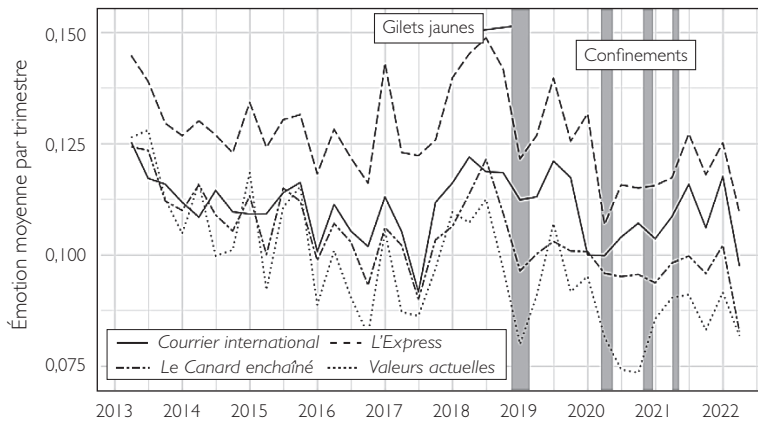


Figure 13 – Hebdomadaires d'information et émotions.

Pour finir ce tour d'horizon du paysage médiatique, nous positionnons les médias web au regard de leurs homologues classiques. Le caractère plus négatif des sentiments exprimés par les abonnés à *Hugo décrypte* et à *Thinker View* paraît structurel durant la période : la moyenne a pratiquement toujours été en dessous de celle correspondant aux titres classiques et se trouve aussi en dessous du niveau d'autres médias web tel que *Les Jours* et *Mediapart* (figure 14).

Notre indice met en évidence une polarisation multiple de l'espace des sentiments. D'un côté, nous retrouvons le contraste politique bien identifié entre le centre et les extrêmes, ces derniers étant caractérisés par un bien-être plus faible et des sentiments plus négatifs. D'un autre côté, et de manière sans doute plus originale, nous mettons en évidence un contraste générationnel, avec des médias jeunes associés à une plus forte expression de sentiments négatifs. Ce mal-être a sans doute trouvé une traduction dans les urnes au premier tour de l'élection présidentielle de 2022, avec un soutien marqué des plus jeunes générations à Jean-Luc Mélenchon.

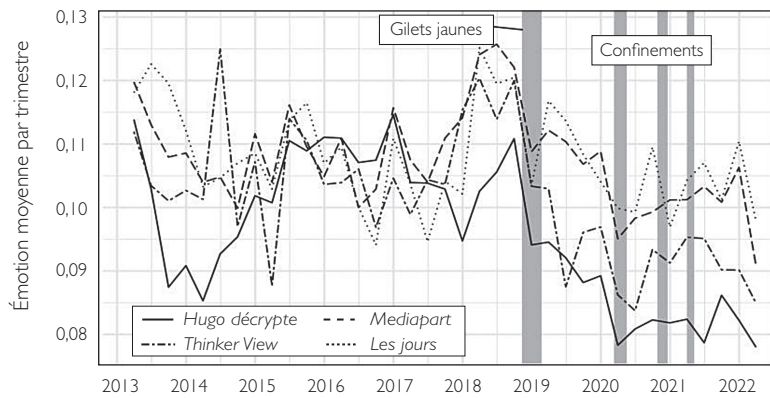


Figure 14 – Sélection de médias.

3. Robustesse et extensions de l'indicateur

COMPARAISON AVEC L'INDICE HEDONOMETER

Afin de mesurer la robustesse de notre mesure, nous comparons notre indicateur à celui calculé à partir de tweets par P. Dodds *et al.*³⁶ et disponible via l'indicateur Hedonometer (figure 15).

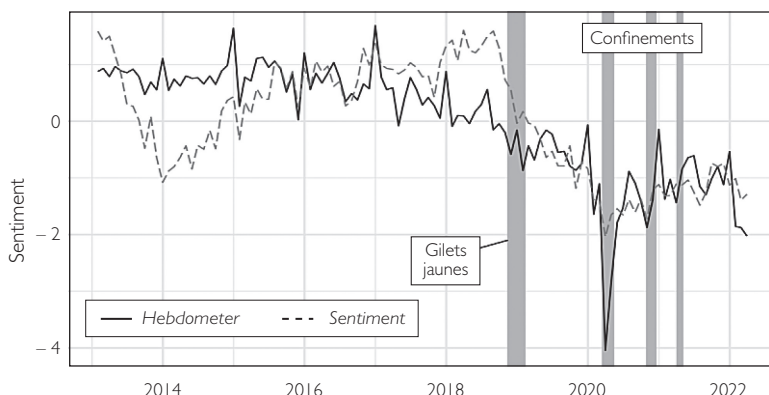


Figure 15 – Indicateur de moral des Français Twitter et Hedonometer.

Notre indicateur de sentiment suit les mêmes mouvements que l'indicateur Hedonometer en français, très utilisé. La corrélation entre les deux indicateurs – à fréquence mensuelle – est de 0,73 pour notre indicateur calculé avec l'approche VADER. Cette analyse permet de confirmer la fiabilité de notre indicateur, construit à partir des données Twitter, en comparaison avec les autres indicateurs existant dans la littérature.

36. Art. cité.

La différence entre notre indicateur de sentiment basé sur VADER et l'indicateur de P. Dodds *et al.*³⁷ basé sur le dictionnaire lab-MT-fr s'explique selon nous par le fait que certains mots sont classifiés de manière trop extrême dans ce dernier dictionnaire. Par exemple, le mot « corona-virus » est le plus négatif de l'ensemble du dictionnaire lab-MT-fr avec un score inférieur à celui de mots comme « suicide », « souffrance » ou « malheur », ce qui nous semble exagéré. La mise à jour du dictionnaire lab-MT-fr en 2020, avec l'inclusion de nouveaux mots liés au coronavirus, entraîne une chute très importante de l'indicateur Hedonometer à partir de 2020 étant donné le volume de tweets contenant le mot « corona-virus » (et ses variantes).

INDICATEUR DE SENTIMENT ET INDICATEURS DE BIEN-ÊTRE

Depuis juin 2016, la plate-forme « Bien-être » de l'enquête mensuelle de conjoncture auprès des ménages (Camme) est administrée trimestriellement, avec vingt questions sur le bien-être subjectif. Une question en particulier porte sur le bien-être émotionnel – « À quel point vous êtes-vous senti(e) heureux(se) hier ? » – et pourrait se rapprocher en particulier de notre indicateur de sentiment. L'enquête étant réalisée entre la fin du mois précédent et le 17 du mois, nous ne retenons pour la comparaison qu'un score de sentiment calculé sur les 15 jours précédant la fin de l'enquête.

L'exercice est structurellement limité par le faible nombre d'observations disponibles. De juin 2016 à décembre 2021, nous ne disposons que de 20 vagues de la plate-forme « Bien-être » de Camme. Il semble toutefois que notre indicateur de sentiment est peu lié aux dimensions du bien-être mesurées par l'enquête. Pratiquement, tous ces indicateurs ont connu un fort rebond à la suite du premier déconfinement, une évolution que nous n'observons pas dans notre indicateur de sentiment

37. Art cité.

via les données Twitter. D'un autre côté, ce dernier indique une dégradation durable depuis 2018 et la crise des Gilets jaunes, alors que celle-ci n'a eu qu'un impact ponctuel sur les indicateurs de bien-être disponibles dans Camme.

Les causes de ces écarts sont potentiellement multiples. L'enquête Camme porte sur un échantillon représentatif de la population française, alors que Twitter constitue un échantillon biaisé. Surtout, nous ne mesurons pas la même chose, en dépit de la ressemblance apparente des indicateurs. Le tweet est le plus souvent une action instantanée, une réaction immédiate à quelque chose qu'on a vu ou lu. On « tweete » parce que l'on ressent une émotion forte et l'envie de l'exprimer. Twitter va donc donner un échantillon de sentiments lorsque les individus sont dans un état émotionnel « chaud », positivement comme négativement. À l'inverse, les questions de la plate-forme « Bien-être » présentent une dimension rétrospective ou réflexive. Lorsque l'on nous demande à quel point nous nous sommes sentis heureux hier, la question invite explicitement à tenir compte de l'ensemble de la journée précédente, non pas seulement des moments de grande joie ou d'énervement. Nous voyons sans doute ici la différence entre une humeur d'ensemble, mesurée par les questions d'enquête, et une humeur plus instantanée avec l'indicateur de sentiment Twitter. D'où l'intérêt de suivre en parallèle ces deux familles d'indicateurs.

QUE MESURE-T-ON EN ANALYSANT LE CONTENU DES MESSAGES ?

L'analyse du contenu des tweets – et plus particulièrement des hashtags – permet de comprendre les principales thématiques ayant entraîné une dégradation du sentiment depuis 2018. Pour cela, nous calculons le sentiment moyen associé à chaque hashtag afin d'identifier pour chaque mois les hashtags associés à un sentiment négatif. Le tableau 3 présente chaque année les vingt hashtags associés à un sentiment négatif.

Tableau 3 – Hashtags les plus fréquemment associés à des émotions négatives (par an)

	date	hashtag
2	2012	['#triste', '#putain', '#beurk', '#merde', '#dilemme', '#dur', '#fuck', '#nul', '#bad', '#geek', '#titanic', '#bizarre', '#vdm', '#csrracing', '#rip', '#desperatehousewives', '#sncl', '#seum', '#cestdit', '#confessionsintimes']
3	2013	['#suicide', '#frip', '#honte', '#stress', '#putain', '#triste', '#marre', '#insomnie', '#malade', '#peur', '#merde', '#fatigue', '#dilemme', '#nul', '#pourquoi', '#splash', '#yenaassez', '#sad', '#greysanatomy', '#dur']
4	2014	['#triste', '#fatigue', '#gameofthrones', '#suivezlemouvement', '#omgserv', '#gameinsi', '#ipa', '#gam', '#sncl', '#csrrclassics', '#seum', '#twd', '#confessionsintimes', '#rip', '#fraall', '#vdm', '#flemme', '#flappybird', '#hollande', '#polqc']
5	2015	['#fusillade', '#trecur', '#sospascal', '#rip', '#facetoface', '#gam', '#clem', '#csrrclassics', '#mtvstars', '#prayforparis', '#game', '#thewalkingdead', '#greysanatomy', '#help', '#hollande', '#got', '#charliehebo', '#harrypotter', '#on', '#vdm']
6	2016	['#prayfornice', '#loitravailnonmerci', '#nice06', '#vdm', '#nice', '#thewalkingdead', '#titanic', '#gam', '#electionnight', '#the100', '#orangesponsorsyou', '#moto', '#twd', '#loitravail', '#brexit', '#héliporteur', '#nikeplus', '#teenwolf', '#resultatsbac', '#oitnb']
7	2017	['#4mariagespour1lunedemi', '#lepen', '#bbmas', '#balancetonporc', '#2017ledebat', '#fcbpsg', '#fn', '#rouen', '#bac2017', '#trump', '#macron', '#fillon', '#lemissionpolitique', '#bourdindirect', '#valls', '#adp2017', '#rdc', '#lrem', '#onpc', '#onelovemanchester']
8	2018	['#giletsjaunes', '#bac2018', '#onpc', '#benalla', '#kohlanta', '#macron', '#adp2018', '#lrem', '#lrt', '#fakenews', '#sncl', '#nowplaying', '#quotidien', '#forzahorizon4', '#xboxpgw', '#eds', '#tpmp', '#pekinexpress', '#teamol', '#rdc']
9	2019	['#notredame', '#giletsjaunes', '#benalla', '#gameofthrones', '#jeudiconfession', '#macron', '#kohlanta', '#lrem', '#rn', '#got', '#lesanges11', '#lrt', '#eurovision', '#polqc', '#balancetonpost', '#thevoice', '#nouvelphotodeprofil', '#', '#europeennes2019', '#france']
10	2020	['#trump', '#covid-19', '#macron20h', '#coronavirus', '#macron', '#lmac', '#covid19france', '#kohlanta2020', '#lrem', '#covid19', '#kohlanta', '#missfrance2021', '#cesoirchezbaba', '#covid_19', '#10edereduction', '#topchef', '#', '#confinement', '#foodballdays', '#nintendoswitch']
11	2021	['#kohlanta2021', '#covid_19', '#covid', '#poupettekenza', '#faceababa', '#macron20h', '#kohlanta', '#vaccin', '#lrem', '#passsanitaire', '#nonaupassdelahonte', '#polqc', '#macron', '#ordm', '#kohlantalalegende', '#tpmp', '#cnews', '#passdelahonte', '#zemmour', '#lrl']
12	2022	['#macrondehors', '#russie', '#poupettekenza', '#poutine', '#ukraine', '#macronnousprendpourdescons', '#lepen', '#cnews', '#toutsaufmacron', '#lfi', '#covid', '#tpmp', '#passvaccinal', '#macrondegage', '#mlp', '#lrl', '#radiolondres', '#presidentielles2022', '#pekinexpress', '#macron']

On retrouve pour chaque année – au milieu d'autres hashtags relatifs à des événements ou à des émissions de télévision – les grandes thématiques ayant impacté négativement le sentiment des Français : les attentats en 2015, la loi Travail en 2016, #balancetonporc en 2017, les Gilets jaunes et l'affaire Benalla en 2018 et 2019, le coronavirus en 2020 et le passe sanitaire en 2021, et #zemmour et #ukraine en 2022. Le hashtag #macron est aussi largement associé à un sentiment négatif qui montre la forte opposition et la polarisation des sentiments autour des actions et de la personnalité du président Macron pendant son mandat (chose que l'on ne retrouvait pas par exemple durant le quinquennat précédent de François Hollande dont le nom n'apparaît pas directement dans la liste des hashtags négatifs).

Nous pouvons aussi mesurer l'évolution de l'intérêt pour différents sujets / thématiques à partir de l'analyse des émojis et des mots-clés utilisés dans les tweets. La figure 16 montre, à fréquence journalière, l'évolution de la fréquence de l'emoji 📄 avec la fréquence des mots-clés « vaccin » et « vaccination » entre 2020 et 2022 (moyenne mobile 30 jours).

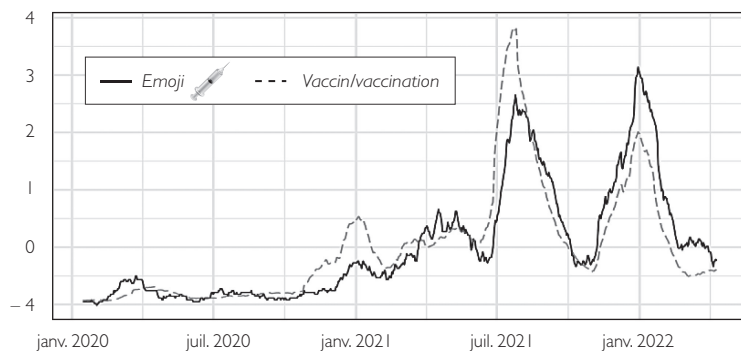


Figure 16 – Évolution du nombre de tweets contenant l'emoji 📄 en 2020-2021.

Cet exemple illustre la possibilité de créer des indicateurs thématiques *via* les données Twitter permettant de suivre de manière plus fine l'évolution des sentiments et des émotions autour d'un sujet spécifique.

LOCALISATION ET ORIGINE DES TWEETS

Afin de déterminer précisément la localisation des messages de notre échantillon et d'identifier les biais géographiques possibles, nous analysons plus en détail l'ensemble des tweets géolocalisés de notre échantillon. Pour ceux-ci – qui représentent environ 1 % de l'ensemble des tweets – nous identifions le pays d'origine ainsi que les coordonnées GPS précises (latitude et longitude). Nous trouvons tout d'abord qu'approximativement 70 % des tweets de notre échantillon proviennent d'utilisateurs localisés en France. Le reste des tweets en français provient du Canada (8 %), du Cameroun (4 %), de Belgique (3 %), de Côte d'Ivoire (2 %), du Sénégal (2 %) et d'autres pays francophones. Cette analyse confirme que nos indicateurs basés sur des tweets en français capturent très majoritairement le sentiment des Français, tout en montrant qu'il peut tout de même exister un biais lié à la présence non négligeable de tweets en provenance de pays francophones. Ce biais peut être diminué lors de l'analyse des tweets des utilisateurs qui suivent les médias français ou des comptes de politiciens français. En ce qui concerne les tweets géolocalisés en France, nous analysons tout d'abord la surreprésentation/sous-représentation de certaines régions, en comparant le nombre de tweets par région au nombre d'habitants de cette même région. Nous trouvons que l'Ile-de-France est surreprésentée et que la Corse est sous-représentée. En ajustant par le nombre d'habitants, il y a environ quatre fois plus d'utilisateurs de Twitter en Ile-de-France qu'en Corse – à supposer que l'utilisation de l'option géolocalisation soit identique partout en France et que le pourcentage des utilisateurs qui « tweetent » depuis leur téléphone soit relativement constant. La figure 17 présente une cartographie réalisée à partir des données Twitter géolocalisées (latitude et longitude).

En dehors de la situation spécifique de Paris, la localisation des tweets est relativement homogène entre les différents départements, ce qui laisse à penser que l'indicateur de moral des Français que nous construisons à partir des données Twitter n'est pas trop biaisé géographiquement.

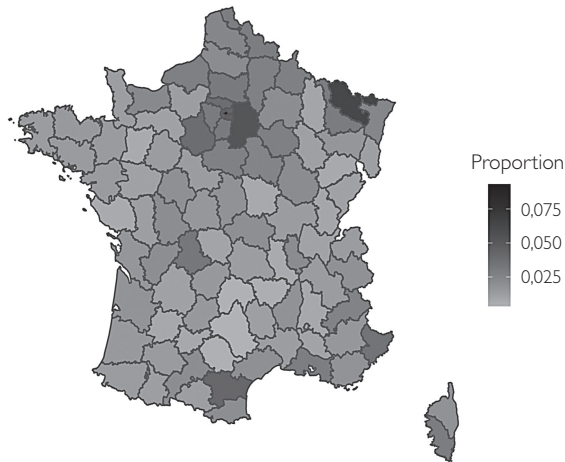


Figure 17 – Localisation des tweets de l'échantillon.

4. Pour ceux qui ont (encore) des doutes sur la validité des données Twitter

MESURER LE BIEN-ÊTRE VIA TWITTER

De nombreux travaux récents ont tenté d'extraire de Twitter des informations sur l'état de l'opinion et sur le bien-être ressenti par la population. On pourra se reporter à la *Note de l'Observatoire du bien-être*, 2020-09, « Twitter, mesure du bien-être ? » pour une revue de la littérature. L'article fondateur de la littérature sur la mesure du bien-être *via* Twitter est celui de P. Dodds *et al.*³⁸. Ses auteurs ont créé un nouvel indicateur – Hedonometer³⁹ – qui mesure l'évolution du sentiment des millions de messages publiés sur Twitter chaque jour. Cet indicateur est mis à jour quasiment en temps réel et l'analyse de sentiment a été étendue depuis à une dizaine de langues dont le français. Cette première étude a permis de mettre en avant des effets de saisonnalité importants : les individus sur Twitter sont plus heureux le week-end que la semaine et plus heureux durant les vacances, avec un pic chaque année autour de Noël et du Nouvel An. Par la suite, à partir du lexique développé par P. Dodds *et al.*⁴⁰, L. Mitchell *et al.*⁴¹ ont utilisé un échantillon de plusieurs millions de tweets géolocalisés afin d'étudier la corrélation – au niveau de chaque État et communauté urbaine aux États-Unis – entre le bien-être mesuré *via* Twitter et le bien-être issu de données de sondages plus traditionnels (entre autres avec le Gallup Sharecare Well-Being Index). Les auteurs montrent qu'il existe une très forte corrélation entre la mesure du bien-être Twitter et les mesures issues de sondage, suggérant

38. Art. cité.

39. <https://hedonometer.org>

40. Art. cité.

41. L. Mitchell, M. Frank, K. Harris, P. Dodds *et* C. Danforth, « The geography of happiness : Connecting twitter sentiment and expression, demographics, and objective characteristics of place », 2013.

donc que cette approche permet d'obtenir des indicateurs robustes en temps réel permettant de suivre l'évolution du bien-être à l'échelle locale ou régionale. À un niveau agrégé, H. Schwartz et *al.*⁴² montrent que l'analyse des tweets provenant de 1 300 comtés américains différents est prédictive du bien-être subjectif des personnes vivant dans ces comtés, tel que mesuré par des enquêtes représentatives. Dans cette étude, bien que les utilisateurs analysés *via* Twitter soient différents des personnes sondées, l'agrégation du bien-être de l'échantillon Twitter est prédictive du bien-être agrégé des sondés. Ces résultats viennent confirmer ceux de D. Quercia et *al.* qui ont montré que le sentiment des individus sur Twitter à Londres et dans la banlieue londonienne était fortement corrélé à l'« Index of Multiple Deprivation » calculé par l'Office statistique national du Royaume-Uni pour 78 zones de recensement⁴³. K. Jaidka et *al.* ont réalisé ce qui nous semble à ce jour l'étude la plus complète et robuste concernant l'utilisation de données Twitter pour la construction d'indicateurs de bien-être subjectif⁴⁴. Les auteurs utilisent un jeu de données de plus d'un milliard de tweets géolocalisés – entre 2009 et 2015 – afin de construire des indicateurs de bien-être au niveau de chaque comté et comparent leurs résultats aux données issues du sondage Gallup-Sharecare Well-Being Index. Ils montrent à ce propos qu'une approche de *machine learning* permet à la fois de mieux capturer le bien-être des individus qu'une approche fondée sur un dictionnaire et d'augmenter significativement la corrélation entre le bien-être mesuré *via* Twitter et le bien-être issu de données d'enquêtes.

En Italie, sous l'impulsion de Andrea Ceron, Luigi Curini, Stefano Iacus et Giuseppe Porro, de nombreuses études ont été réalisées utilisant des

42. H. Schwartz, J. Eichstaedt, M. Kern, L. Dziurzynski, R. Lucas, M. Agrawal et L. Ungar, « Characterizing geographic variation in well-being using tweets », 2013.

43. D. Quercia, J. Ellis, L. Capra et J. Crowcroft, « Tracking gross community happiness from tweets », 2012.

44. Art. cité.

tweets géolocalisés en italien afin de dériver des mesures du bien-être à l'échelle régionale et nationale (L. Curini et *al.*⁴⁵, S. Iacus et *al.*⁴⁶). Les auteurs ont proposé une méthodologie permettant de construire un indice de bien-être à partir du contenu publié sur les réseaux sociaux – le Social Well-Being Index ou SWBI⁴⁷. S. Iacus et G. Porro ont d'ailleurs publié en mai 2021 un ouvrage complet de 220 pages, intitulé *Subjective Well-Being and Social Media*, qui synthétise de manière remarquable l'ensemble de la littérature sur la mesure du bien-être *via* Twitter. En Turquie, A. Durahim et M. Coskun ont construit un indicateur de Gross National Happiness à partir des données Twitter et ont comparé leurs résultats à ceux d'un indicateur de bien-être produit par l'Institut national statistique turc à partir de sondages et disponible pour chaque localité⁴⁸. Les auteurs trouvent une corrélation faible entre la mesure Twitter et la mesure *via* sondage et attribuent cette différence au fait que la mesure *via* sondage est un « instantané » à une date donnée, tandis que leur indicateur Twitter capture un sentiment sur l'ensemble d'une période. En Chine, où Twitter est interdit, B. Hao et *al.* ont étudié les données du réseau social Weibo et réalisé une expérimentation en demandant à des individus de remplir un questionnaire psychologique, avant de comparer les résultats du questionnaire au sentiment exprimé par ces mêmes individus sur Weibo⁴⁹.

45. L. Canova, L. Curini et S. Iacus, « Measuring idiosyncratic happiness through the analysis of Twitter : An application to the Italian case », 2014.

46. S. Iacus, G. Porro, S. Salini et E. Siletti, « An Italian composite subjective well-being index : The voice of Twitter users from 2012 to 2017 », 2020.

47. Récemment, un indicateur SWBI a d'ailleurs été calculé pour le Japon par T. Carpi, A. Huno, S. Iacus et G. Porro, « Twitter subjective well-being indicator during Covid-19 Pandemic : A cross-country comparative study », *arXiv*, 19 janvier 2021.

48. A. Durahim et M. Coşkun, « #iamhappybecause : Gross national happiness through Twitter analysis and big data », 2015.

49. B. Hao, L. Li, R. Gao, A. Li et T. Zhu, « Sensing subjective well-being from social media », 2014.

Récemment, à la suite de l'ouverture d'un jeu de données Twitter durant la pandémie, de nombreuses recherches ont analysé l'impact de la Covid-19 sur le sentiment et les émotions des individus (M. Lwin et al.⁵⁰ ; J. Xue et al.⁵¹). Pour leur part, Y. Wang et al. ont analysé les données de Weibo pour mesurer les différences de bien-être subjectif des confinements dans différentes villes chinoises⁵². F. Sarracino et al. ont montré que durant la pandémie de Covid, un indicateur construit à partir de Twitter pour dix pays produisait des résultats comparables à ceux de l'Eurobaromètre et du *World Happiness Report*⁵³. Ces auteurs publient leur indicateur (Gross National Happiness Today) avec une fréquence journalière depuis 2020⁵⁴. F. Díaz et P. Henríquez ont également eu recours aux données issues de Twitter (et Google Trends) pour mesurer les effets sur le bien-être des confinements différenciés mis en place dans les municipalités chiliennes⁵⁵. Utilisant la forte hétérogénéité sociale de ces municipalités, ils ont pu montrer que le confinement avait eu un effet négatif sur le bien-être des habitants des villes les plus aisées.

50. M. Lwin, J. Lu, A. Sheldenkar, P. Schulz, W. Shin, R. Gupta et Y. Yang, « Global sentiments surrounding the Covid-19 pandemic on Twitter : Analysis of Twitter trends », 2020.

51. J. Xue, J. Chen, R. Hu, C. Chen, C. Zheng, Y. Su, T. Zhu, « Twitter discussions and emotions about the Covid-19 pandemic : Machine learning approach », 2020.

52. Y. Wang, P. Wu, X. Liu, S. Li, T. Zhu et N. Zhao, « Subjective well-being of Chinese Sina Weibo users in residential lockdown during the Covid-19 Pandemic : Machine learning analysis », 2020.

53. F. Sarracino, T. Greyling, K. O'Connor, C. Peroni et S. Rossouw, « Trust predicts compliance to Covid-19 containment policies : Evidence from ten countries using big data », 2021.

54. <http://gnh.today>

55. F. Díaz et P. Henríquez, « Social sentiment segregation : Evidence from Twitter and Google trends in Chile during the Covid-19 dynamic quarantine strategy », 2021.

Dans un chapitre du *World Happiness Report*, C. Bellet et P. Frijters reprennent différentes études afin de discuter de la prédictibilité de la satisfaction dans la vie au niveau individuel, local et national *via* l'utilisation des données *big data* (issus de Twitter, mais aussi de Google Trends et Facebook, voir Annexe 2)⁵⁶. La synthèse des résultats révèle la très forte hétérogénéité de la prédictibilité du bien-être avec des données *big data*. En fonction de la période d'étude, de la fréquence d'analyse, du pays considéré, de la mesure de bien-être (satisfaction dans la vie ou bonheur) et de l'unité d'analyse (individuel, local, agrégé), les résultats varient fortement d'une étude à l'autre. Dans le dernier rapport du *World Happiness Report*, H. Metzler *et al.* s'intéressent spécifiquement à la mesure du bien-être *via* Twitter lors de la pandémie de la Covid⁵⁷. Les auteurs confirment, à travers leurs mesures, que les sentiments d'anxiété ou de tristesse ont augmenté durant les premières semaines de la période Covid dans les dix-huit pays qu'ils ont considérés, avec un effet légèrement plus important pour les femmes que pour les hommes.

Pour reprendre la classification de Diener – classification reprise par exemple par le *World Happiness Report* – le bien-être comporte trois principales dimensions que sont la satisfaction dans la vie, les émotions positives et les émotions négatives⁵⁸. La satisfaction dans la vie est une mesure relativement stable tandis que les émotions (les affects) positives et négatives sont plus changeantes en fonction des événements. À partir d'une même source de données, il est possible, en choisissant certains mots-clés ou en utilisant un dictionnaire spécifique, de construire des proxys de la satisfaction dans la vie ou des proxys des affects. F. C. Yang et P. Srinivasan se concentrent uniquement sur les tweets contenant explicitement des mots en lien avec la satisfaction dans la vie – par exemple

56. C. Bellet et P. Frijters, « Big data and well-being », 2020.

57. H. Metzler, M. Pellert et D. Garcia, « Using social media data to capture emotions before and during Covid-19 », 2022.

58. E. Diener, « Subjective well-being », *Psychological Bulletin*, 95 (3), 1984.

« je déteste / j'aime ma vie » – pour constituer leur échantillon⁵⁹. Par la suite, ils interrogent directement un échantillon d'utilisateurs Twitter afin de comparer la mesure calculée à partir des messages de ces utilisateurs et leur bien-être subjectif exprimé *via* le sondage ; ils trouvent une corrélation élevée entre la satisfaction déclarée sur Twitter et la satisfaction exprimée *via* sondage. Cela tend à confirmer la validité de l'utilisation des tweets comme indicateur de la satisfaction dans la vie, tout du moins lorsqu'un filtre strict est appliqué sur les messages afin de ne garder que les tweets qui évoquent directement cette notion de satisfaction dans la vie. Cela implique de ne garder qu'un très faible échantillon de l'ensemble des tweets – environ 0,1 % dans l'échantillon de F. C. Yang et P. Srinivasan.

LES LIMITES DES DONNÉES EN LIGNE

Pour reprendre le titre d'un article de A. Waldherr *et al.*, on peut affirmer que le « *big data comes with big noise* » (« Le big data s'accompagne de beaucoup de bruits »)⁶⁰. En ce qui concerne la mesure du bien-être *via* les données Twitter, les biais peuvent principalement être causés par la non-représentativité des utilisateurs Twitter, le faible ratio « signal-bruit » et le risque de manipulation de ces indicateurs.

Premièrement, la non-représentativité des utilisateurs Twitter est sûrement le biais le plus communément cité dans les études utilisant des données *big data* pour mesurer le bien-être subjectif. Selon une étude du Pew Research Center réalisée en 2019, les utilisateurs Twitter – tout du moins aux États-Unis – sont plus jeunes, plus éduqués et plus urbains que la population générale⁶¹. Des résultats similaires ont été mis en avant par

59. Art. cité.

60. A. Waldherr, D. Maier, P. Miltner et E. Günther, « Big data, big noise : The challenge of finding issue networks on the web », 2017.

61. <https://www.pewresearch.org/internet/2019/04/24/sizing-up-twitter-users/>

A. Mislove *et al.*⁶². Dans une seconde étude, le Pew Research Center a montré en 2021 que les 18-29 ans sont plus susceptibles de se connecter sur la plateforme et d'utiliser Twitter comme une source d'information, ce qui permet notamment d'améliorer la compréhension de certains événements (notamment en direct) et de favoriser l'engagement politique⁶³. Un autre élément particulièrement important souligné dans cet article est le fait que sur Twitter, 97 % des tweets sont produits par 25 % des utilisateurs. On peut ainsi distinguer deux types d'utilisateurs sur Twitter, ceux qui ont une utilisation active, à travers leurs messages, leurs *retweets*, et ceux qui ont une utilisation plus passive, mais qui peuvent néanmoins, à travers des *retweets* ou des *likes*, nous fournir des informations sur leurs opinions ou leurs sentiments.

K. Jaidka *et al.* montrent que l'échantillon des sondés du Gallup Well-Being Index est plus âgé et plus aisé que la « population Twitter ». Ce biais tend d'ailleurs à être renforcé dans les études concentrées sur les tweets géolocalisés qui proviennent des messages envoyés depuis l'application mobile de Twitter et pour lesquels la population est différente de la population générale Twitter (plus jeune encore). Pour tenir compte de la non-représentativité des utilisateurs sur Twitter, les auteurs proposent une post-stratification des échantillons de Gallup et de Twitter de façon à ce qu'ils correspondent aux données démographiques du recensement en termes d'âge, de sexe, d'éducation et de revenus. Ces variables n'étant pas disponibles directement sur les comptes Twitter des utilisateurs, K. Jaidka *et al.* utilisent des méthodes d'analyses textuelles afin d'inférer ces caractéristiques à partir du contenu des messages des utilisateurs⁶⁴. Pour cela, les auteurs utilisent

62. A. Mislove, S. Lehmann, Y. Y. Ahn, J.P. Onnela et N. Rosenquist, « Understanding the demographics of Twitter users », 2011.

63. <https://www.pewresearch.org/internet/2021/11/15/the-behaviors-and-attitudes-of-u-s-adults-on-twitter/>

64. Art. cité.

le lexique développé par M. Sap et al.⁶⁵ et H. Schwartz et al.⁶⁶ à partir de statuts Facebook croisés avec des données sociodémographiques collectées *via* l'application Facebook de test de personnalité « myPersonality »⁶⁷. S. Giorgi et al. ont montré par la suite que la correction des biais de sélection sociodémographiques pour la prédiction des populations à partir des médias sociaux pouvait permettre d'améliorer significativement la précision des indicateurs de satisfaction dans la vie issus de Twitter⁶⁸. Sur les données italiennes, S. Iacus et al.⁶⁹ proposent une méthode de correction du biais d'échantillonnage considérant le taux d'utilisation d'internet et de Twitter dans chaque province d'Italie, afin de prendre en compte la plus faible utilisation de Twitter dans les provinces rurales⁷⁰. Ainsi, en Italie, 92 % des habitants ont un accès Internet et 59 % sont des utilisateurs actifs des réseaux sociaux ; mais ces chiffres varient fortement d'une région à l'autre. Cette méthode implique de travailler sur des messages géolocalisés – qui ne représentent que 1 % des messages publiés sur Twitter – ou bien de dériver la localisation à partir de la description de chaque utilisateur (le champ

65. M. Sap, G. Park, J. Eichstaedt, M. Kern, D. Stillwell, M. Kosinski, L. Ungar et H. Schwartz, « Developing age and gender predictive lexica over social media », 2014.

66. H. Schwartz, J. Eichstaedt, M. Kern, L. Dziurzynski, S. Ramones, M. Agrawal et L. Ungar, « Personality, gender, and age in the language of social media : The open-vocabulary approach », 2013.

67. M. Kosinski, D. Stillwell et T. Graepel, « Private traits and attributes are predictable from digital records of human behavior », 2013.

68. S. Giorgi, V. Lynn, S. Matz, L. Ungar et A. Schwartz, « Correcting sociodemographic selection biases for accurate population prediction from social media », 2021.

69. Art. cité.

70. Même dans les pays développés, environ 10 % des individus n'ont pas d'accès à Internet selon les chiffres de l'International Telecommunication Union. L'importance du redressement lié au pourcentage d'utilisation d'Internet dépend de chaque pays.

« Localisation » sur Twitter est un champ libre et un utilisateur est libre d'indiquer ou non sa localisation).

Deuxièmement, en ce qui concerne le faible ratio signal-bruit, un pourcentage important de tweets sont envoyés par des robots, ce qui risque de biaiser les indicateurs de bien-être subjectif et ce, d'autant plus en cas de construction d'indicateurs agrégés sans ajustement en fonction du volume de messages envoyés par chaque utilisateur – les *bots* étant très actifs sur Twitter. Il existe cependant de nombreuses méthodes pour détecter ces comptes. C. Davis *et al.* proposent un système de détection des robots à partir de l'analyse des tweets (contenu, fréquence, régularité), du réseau (amis et *followers*) et de centaines d'autres variables explicatives⁷¹. Les auteurs ont développé une API gratuite et un site web permettant de calculer un score Botometer – entre 0 et 5 – pour chaque compte (voir F. C. Yang *et al.*⁷² pour une explication détaillée du fonctionnement de Botometer; et E. Alothali *et al.*⁷³, pour une revue de la littérature sur la détection des *bots* sur Twitter). Pour améliorer le ratio « signal-bruit », on peut aussi filtrer les tweets en appliquant une suite de règles ou bien en restreignant l'analyse à certaines thématiques. Les sujets discutés par les utilisateurs sur Twitter sont très variés – sport, jeux vidéo, politique, news, vie personnelle, vie professionnelle – et il est possible que seuls certains d'entre eux soient pertinents pour mesurer le bien-être. F. C. Yang et P. Srinivasan⁷⁴ appliquent une méthode de modélisation de thématique (allocation de Dirichlet latente) pour identifier automatiquement la thématique de chaque message avant de réaliser des indicateurs de bien-être thématique.

71. C. Davis, O. Varol, E. Ferrara, A. Flammini et F. Menczer, « Botornot : A system to evaluate social bots », 2016.

72. F. C. Yang, E. Ferrara et F. Menczer, « Botometer 101 : Social bot practicum for computational social scientists », 2022.

73. E. Alothali, N. Zaki, E. Mohamed et H. Alashwal, « Detecting social bots on Twitter : A literature review », 2018.

74. Art. cité.

Enfin, et troisièmement, un risque commun à l'utilisation de ces nouvelles données est le risque de manipulation des indicateurs. D. Lazer soulignait déjà ce problème pour la prévision des épidémies de grippe, en indiquant que les données Twitter peuvent être manipulées à la fois par le fournisseur de services et par les utilisateurs (comme les entreprises commercialisant un produit)⁷⁵. Le risque le plus important est que des individus tentent de manipuler le processus de génération de données pour atteindre leurs propres objectifs (gain économique ou politique par exemple). Les fausses informations sont déjà très nombreuses sur les réseaux sociaux – comme cela a été particulièrement visible au moment des élections américaines⁷⁶ – et le risque de manipulation est d'autant plus grand que les indicateurs seront ensuite utilisés dans des décisions de politiques publiques.

Plus généralement, les résultats issus des données *big data* ne sont donc pas un substitut aux données de sondage, mais bien un complément⁷⁷. L'analyse *big data* nécessite de construire soigneusement un cadre de recherche, afin à la fois de limiter au maximum les risques de manipulation et d'augmenter le ratio signal-bruit⁷⁸. De plus, à long-terme, des questions relatives à l'endogénéité peuvent se poser : les recherches sur l'utilisation de Twitter pour mesurer le bien-être peuvent pâtir d'un biais lié à l'impact négatif causal de l'utilisation des réseaux sociaux sur le bien-être (voir Annexe 3). Cependant, tout en reconnaissant leurs limites, nous restons convaincus de la possibilité de construire *via* les données Twitter des indicateurs pertinents et robustes permettant de mesurer les sentiments et les émotions des Français en complément des approches traditionnelles et des indicateurs existants.

75. D. Lazer, R. Kennedy, G. King et A. Vespignani, « The parable of Google flu : traps in big data analysis », 2014.

76. H. Allcott et M. Gentzkow, « Social media and fake news in the 2016 election », 2017.

77. C'est le point de vue défendu dans L. Japek et. al, « Big data in survey research », AAPOR Task force report, 2015.

78. L. Einav et J. Levin, « Economics in the age of big data », 2014.

Conclusion

Les mesures fondées sur les données en ligne offrent aux chercheurs en sciences sociales de nombreuses nouvelles opportunités : mesure quasiment en temps réel et à haute-fréquence, analyse détaillée des émotions, analyse des thématiques impactant le bien-être ou encore désagrégation selon des caractéristiques sociodémographiques. Les difficultés posées par ce type de sources sont multiples (biais de représentativité, désirabilité sociale, dynamique spécifique à une plateforme, etc.), mais des travaux récents proposent des solutions à ces différents problèmes. Au total, malgré leurs limites, les données Twitter permettent de prendre de manière fiable le « pouls de la nation », de quantifier la teneur des messages que les Français ont voulu exprimer et d'être utiles ainsi aux décideurs publics.

La dégradation du score de sentiment à partir de l'automne 2018, prélude à la crise des Gilets jaunes, est un fait saillant qui ressort de l'ensemble de nos analyses. Cette dégradation s'est accompagnée d'une polarisation de plus en plus marquée entre les sentiments et les émotions des abonnés à des comptes d'extrême-droite et d'extrême-gauche, d'une part, et ceux des abonnés à des comptes plus proches des partis traditionnels, d'autre part. L'analyse des messages émis par les abonnés aux comptes des personnalités politiques et des médias permet d'illustrer d'une nouvelle manière – et de quantifier – cette polarisation du paysage politique français.

Les données à haute-fréquence permettent également de mesurer l'importance des chocs et la longueur de leur écho dans les messages des Français. Les attentats (Bataclan, *Charlie Hebdo*, Nice), les annonces de confinement, la guerre en Ukraine et la mort soudaine d'une célébrité (Kobe Bryant) sont les événements qui ont le plus dégradé le moral des Français au cours de la période étudiée. Mais certains chocs n'ont qu'un impact limité dans le temps (mort d'une célébrité) tandis que d'autres

— comme le confinement — exercent une influence beaucoup plus durable sur le moral des Français.

Au total, en permettant de quantifier la tonalité des émotions exprimées sur Twitter — de manière agrégée ou en fonction des préférences politiques, de l'âge ou du genre —, ces indicateurs proposent un nouveau regard sur les facteurs influençant le bonheur des Français. Afin d'encourager le développement de tels travaux, l'ensemble des séries élaborées sont disponibles en ligne sur le site de l'Observatoire du bien-être.

Annexe 1. Analyse de sentiment

Le tableau A1 – repris de A. Reagan et *al.*⁷⁹ et complété dans le cadre de cette étude – présente une liste de dictionnaires existants et les principales approches utilisées pour créer chaque dictionnaire. Tous les dictionnaires sont disponibles en ligne gratuitement pour les chercheurs, à l'exception du LIWC – Linguistic Inquiry and Word Count. Comme le montre ce tableau, certains dictionnaires utilisent une approche binaire – c'est-à-dire qu'un score de – 1 est attribué à chaque mot défini comme négatif et un score de + 1 à chaque mot défini comme positif ; d'autres utilisent une échelle plus large permettant de différencier par exemple les mots « légèrement négatifs » des mots « très négatifs ». On observe aussi des différences en ce qui concerne la méthodologie utilisée afin de construire les dictionnaires, avec principalement deux approches : une approche purement manuelle où des experts vont indiquer une polarité pour chaque mot en fonction de leur opinion ; et une approche « crowd-sourcée » – dans laquelle il va être demandé à des sondés d'indiquer leur perception de la polarité de chaque mot – combinée parfois à une utilisation d'Amazon Mechanical Turk (voir encadré I).

**Tableau A1 – Dictionnaire existant pour l'analyse de sentiment
(tiré de A. Reagan et *al.*, 2017)**

Dictionnaire	#Mots	Score	Construction	Source
labMT	10 222	1.3 → 8.5	Sondage + Amazon Mechanical Turk	Dodds P. et al. (2011)
ANEW	1 034	1.2 → 8.8	Survey	Bradley M.M., et Lang P.J. (1999)
LIWC07	2 145	[-1, 0, 1]	Manuelle	Pennebaker J.W., Francis M.E. et Booth R.J. (2001)

79. A. Reagan, C. Danforth, B. Tivnan, J. R. Williams et P. Dodds, « Sentiment analysis methods for understanding large-scale texts : A case for using continuum-scored words and word shift graphs », 2017.

Dictionnaire	#Mots	Score	Construction	Source
MPQA	5 587	[-1, 0, 1]	Manuelle et <i>Machine Learning</i>	Wilson T., Wiebe J. et Hoffmann P. (2005)
WK	13 915	1.3 → 8.5	Sondage : Mechanical Turk	Warriner A.B., Kuperman V. et Brysbaert M. (2013)
LIWC01	2 322	[-1, 0, 1]	Manuelle	Pennebaker J.W., Francis M.E. et Booth R.J. (2001)
LIWC15	6 549	[-1, 0, 1]	Manuelle	Pennebaker J.W., Francis M.E. et Booth R.J. (2001)
PANAS-X	20	[-1, 1]	Manuelle	Watson D. et Clark L.A. (1999)
Pattern	1 528	[-5, -4, , 4, 5]	Non spécifié	De Smedt T. et Daelemans W. (2012)
SentiWordNet	147 700	[-1, 1]	Synset et synonymes	Baccianella S., Esuli A. et Sebastiani F. (2010)
AFINN	2 477	0.0 → 3.0	Manuelle	Nielsen F. (2011)
GI	3 629	[-1, 0, 1]	Harvard-IV-4	Stone P.J., Dunphy D.C. et Smith M.S. (1966)
WDAL	8 743	-1.0 → 1.0	Sondage auprès d'étudiants	Whissell C. <i>et al.</i> (1986)
EmoLex	14 182	-6.9 → 7.5	Sondage : Mechanical Turk	Mohammad S.M. et Turney P.D. (2013)
MaxDiff	1 515	-5.0 → 5.0	Sondage : Mechanical Turk, MaxDiff	Kiritchenko S, Zhu X et Mohammad S.M. (2014)
HashtagSent	54 129	-30.2 → 30.7	Sondage	Zhu X., Kiritchenko S. et Mohammad S.M. (2014)
Sent140Lex	62 468	-1.0 → 1.0	Sondage	Mohammad S.M., Kiritchenko S. et Zhu X. (2013)
SOCAL	7 494	[-5, -4, , 4, 5]	Manuelle	Taboada M. <i>et al.</i> (2011)
SenticNet	30 000	-1.0 → 1.0	Label propagation	Cambria E. <i>et al.</i> (2014)

Dictionnaire	#Mots	Score	Construction	Source
Emoticons	132	[-1, 0, 1]	Manuelle	Gonçalves P. <i>et al.</i> (2013)
SentiStrength	1 270	[-5, -4, , 4, 5]	LIWC + GI	Thelwall M. <i>et al.</i> (2010)
VADER	7 502	-3.9 → 3.4	Sondage Mechanical Turk	Hutto C.J. et Gilbert E. (2014)
Umigon	927	[-1, 1]	Manuelle	Levallois C. (2013)
USent	592	[-1, 1]	Manuelle	Pappas N., Katsimpras G. et Stamatatos E. (2013)
EmoSenticNet	13 188	[-10, -2, -1, 0, 1, 10]	Sondage	Poria S., Gelbukh A. <i>et al.</i> (2013)

Encadré 1 – Amazon Mechanical Turk et la construction de dictionnaire de sentiment

La construction d'un dictionnaire de mots pour l'analyse de sentiment est une tâche longue – qui implique d'analyser et de donner un score à chaque mot – et sujette à un biais de la part de l'annotateur. Afin de pallier ces deux problèmes, il est possible d'utiliser un système d'annotation externalisé consistant à demander à un grand nombre d'individus de classer chaque mot-clé en échange d'une rémunération pour chaque micro-tâche. Le concept de micro-tâche a été popularisé par Amazon *via* le lancement en 2005 de la plateforme Amazon Mechanical Turk. Cette plateforme permet à des *turkers* – le nom des travailleurs sur cette plateforme – de répondre aux demandes des *requesters* – le nom des personnes proposant des tâches à réaliser en échange d'une rémunération. Il y aurait en 2021 environ 30 000 *turkers* sur la plateforme, dont le salaire horaire moyen est de 2 dollars.

Comme discuté par Sorokin et Forsyth⁸⁰, Amazon Mechanical Turk est fréquemment utilisé pour le travail d'annotation de texte ou d'image dans le but de construire ensuite des dictionnaires ou bien d'entraîner des algorithmes de *machine learning*. Une micro-tâche peut consister à indiquer pour un message donné le sentiment associé (négatif/neutre/positif) ou l'émotion principale (joie, tristesse, colère, peur). Pour vérifier la qualité des données, il peut être demandé à plusieurs annotateurs de classer certains messages afin de calculer la similarité entre les classifications.

Amazon Mechanical Turk a été utilisé par P. Dodds *et al.* afin d'associer un score entre 1 et 9 – de négatif à positif – aux mots-clés les plus fréquemment utilisés dans une dizaine de langues⁸¹. Chaque mot a été classifié par 50 annotateurs sur Amazon Mechanical Turk, ce qui permet, en plus de la moyenne du score de sentiment, de calculer l'écart-type des notations afin d'identifier les mots pour lesquels le consensus est plus faible. P. Dodds *et al.* recommandent de supprimer les mots dont la moyenne des sentiments est trop proche de 5, ces mots étant principalement des mots neutres et des *stopwords* qui peuvent biaiser la classification.

Parmi l'ensemble de ces dictionnaires, certains ont été créés plus particulièrement pour l'analyse du sentiment des messages publiés sur les réseaux sociaux. Les messages sur Twitter – comparés par exemple à des articles de médias – sont en effet plus courts, informels et utilisent un langage spécifique composé entre autres d'emojis et de hashtags. Ces dictionnaires sont labMT développé par P. Dodds *et al.*⁸² ; VADER

80. A. Sorokin et D. Forsyth, « Utility data annotation with Amazon Mechanical Turk », *Computer Vision and Pattern Recognition Workshops*, 2008.

81. Art. cité.

82. Art. cité.

développé par C. Hutto et al.⁸³ ; Umigon développé par Levallois⁸⁴ ; NRC MaxDiff⁸⁵ ; NRC Hashtag Sent et NRCSent140 du National Research Council du Canada. Le choix du dictionnaire dépend donc du contexte, du type de texte à analyser et du langage. Il s'agit d'une question empirique qui impose donc de disposer d'un jeu de données de test – pour lequel le sentiment de chaque tweet est connu – ce qui permet d'évaluer la précision des différentes approches. En anglais, ces jeux de données sont nombreux – voir par exemple « STS Gold Dataset »⁸⁶ ou « SemEval-2017 »⁸⁷ – et sont largement utilisés pour comparer la précision des différents dictionnaires et approches. Comme montré par A. Reagan et al., les dictionnaires développés spécifiquement pour l'analyse de tweets ont une précision de classification hors échantillon (sur le jeu de données STS Gold Dataset) bien supérieure à la précision des dictionnaires traditionnels⁸⁸.

La majorité des études sur l'analyse de sentiment sont réalisées à partir de tweets en anglais, mais il existe des versions non-anglophones de certains dictionnaires. Ainsi, la liste de mots-clés de P. Dodds et al.⁸⁹ est disponible en français sur le site Hedonometer, et une version française du package VADER – intitulé VADERsentiment-FR – est disponible sur GitHub. D'autres approches sont possibles, ainsi la traduction en français des dictionnaires en anglais, spécifique à l'analyse de sentiment,

83. Art. cité.

84. C. Levallois, « Umigon : Sentiment analysis for tweets based on lexicons and heuristics », *7th International Workshop on Semantic Evaluation*, Atlanta, 2013.

85. S. Kiritchenko, X. Zu et S. Mohammad, « Sentiment analysis of short informal texts », *Journal of Artificial Intelligence Research (JAIR)*, 24 août 2014.

86. H. Saif, M. Fernandez, Y. He et H. Alani, « Evaluation datasets for Twitter sentiment analysis : A survey and a new dataset, the STS-Gold », *Approaches and perspectives from AI (ESSEM) at AI*IA Conference*, Turin, 2013.

87. S. Rosenthal, N. Farra et P. Nakov, « SemEval-2017 task 4 : Sentiment analysis in Twitter », *Proceedings of the 11th International Workshop on Semantic Evaluation*, Vancouver, 2017.

88. Art. cité.

89. Art. cité.

mais cette méthode tend à faire perdre le contexte et ne permet pas de capturer toutes les spécificités de chaque langue⁹⁰.

En parallèle des méthodes d'analyse de sentiment à partir de dictionnaire – qui ont l'avantage d'être simples et transparentes – il est possible de mettre en œuvre des méthodes de classification automatique supervisée – *machine learning* – afin d'identifier un score de sentiment pour chaque message. Pour utiliser un algorithme de classification supervisée, il faut disposer d'un jeu de tweets classifiés afin d'entraîner un algorithme à identifier les mots et les structures permettant de déterminer la polarité d'un message. La constitution d'un jeu d'entraînement peut se faire manuellement – en demandant à une ou plusieurs personnes de classer chaque message – ou bien semi-automatiquement, en utilisant par exemple les émoticônes/emojis afin de déterminer automatiquement la polarité d'un message⁹¹. La première méthode – *via* la constitution d'un jeu de données d'entraînement classifié manuellement – est coûteuse car elle impose de disposer de plusieurs dizaines de milliers de messages classifiés pour espérer obtenir une précision *out-of-sample* (hors échantillon) suffisante⁹². La seconde méthode – fondée sur l'utilisation des emojis – peut être mise en place à moindre coût mais peut produire d'autres biais, les emojis étant utilisés majoritairement par les jeunes et par les utilisateurs tweetant depuis leur téléphone portable.

90. M. Kaity et V. Balakrishnan, « An automatic non-English sentiment lexicon builder using unannotated corpus », *The Journal of Supercomputing*, 75 (4), 2019.

91. A. Go, R. Bhayani et L. Huang, « Twitter sentiment classification using distant supervision », *Project report, Stanford*, 1 (12), décembre 2009.

92. Sur un jeu de données de messages publiés sur un réseau social dédié à la finance, T. Renault (2020) a montré que la précision de la classification augmentait avec la taille de l'échantillon, mais qu'un plateau semblait être atteint autour de 250 000 tweets, l'amélioration de la précision étant marginalement très décroissante au-delà de ce niveau. Ranco *et al.* (2015) utilisent quant à eux un jeu de données de plus de 100 000 tweets classifiés pour entraîner un algorithme de *machine learning* pour l'analyse de sentiment.

Les approches discutées jusqu'à présent calculent un score de sentiment mais ne prennent pas en compte les différentes émotions sous-jacentes (peur, joie, tristesse, dégoût, colère, surprise). De nombreux auteurs, tels F. Sarracino *et al.*, ont cependant souligné l'importance de la prise en compte des émotions – encore plus dans le cadre de la mesure du bien-être à partir de données textuelles⁹³. Tout comme pour la mesure du sentiment, il est possible de quantifier les émotions en utilisant des dictionnaires existants ou des méthodes de *machine learning*. Le dictionnaire le plus utilisé pour l'analyse d'émotion est le NRC Emotion Lexicon de Saif *et al.*⁹⁴, utilisé par exemple par Greyling *et al.*⁹⁵ (2020) pour identifier les tweets contenant des mots relatifs aux 8 émotions de la classification de Plutchik⁹⁶ et pour calculer des indicateurs de bien-être. Une traduction de ce dictionnaire en français – et un enrichissement *via* les synonymes et les antonymes de chaque mot – a été réalisée par A. Abdaoui *et al.* *via* le dictionnaire « French Expanded Lexicon » (FEEL) qui contient une liste de mots relatifs aux six émotions : joie, colère, tristesse, peur, dégoût, surprise⁹⁷. Tout comme pour l'analyse de sentiment, il existe des jeux de données permettant de tester empiriquement la précision des différentes approches de classification des émotions (mais uniquement en anglais). Pour l'analyse d'émotion sur Twitter, les deux principaux jeux de données sont le « SemEval-2018 Task I : Affect in Tweets » et le « WASSA 2017 », qui associent à chaque tweet un score en fonction de quatre émotions : la peur, la tristesse, la joie et la colère.

Une dernière approche, qui présente l'avantage de ne pas nécessiter un dictionnaire spécifique pour chaque langue, consiste à se concentrer

93. Art. cité.

94. Art. cité.

95. Greyling *et al.* (2020)

96. R. Plutchik, « A general psychoevolutionary theory of emotion », in *Theories of emotion*, Academic Press, 1980.

97. A. Abdaoui, J. Azé, J., S. Bringay et P. Poncelet, « FEEL : A French expanded emotion lexicon », *Springer Verlag*, 51 (3), 2017.

sur les emojis et les émoticônes qui représentent un langage universel et qui sont largement utilisés sur les réseaux sociaux. Par exemple, un message contenant l'emoji 😊 aura tendance à être positif – peu important la langue – et un message contenant l'emoji 😞 aura tendance à être négatif. À partir de l'analyse de plusieurs centaines de milliers de tweets classifiés dans 13 langues, P. Kraji Novak et al.⁹⁸ ont calculé un score de sentiment pour chaque emoji et publié un « Emoji Sentiment Ranking v1.0 » en ligne⁹⁹. Le tableau 2, reconstruit à partir de la base de données de P. Kraji Novak et al., montre le score de sentiment – entre [-1 ; +1] – pour une liste d'emojis. Les auteurs ont montré de plus que le score de sentiment était similaire selon la langue considérée, ce qui tend à montrer que les emojis constituent bien une manière universelle de révéler son ressenti. Le travail a été réalisé ici pour le sentiment, mais il est aussi possible d'associer des emojis à des émotions. C'est l'approche utilisée par A. Shoeb et G. De Melo qui – à partir d'un jeu de données de messages annotés en fonction de différentes émotions – ont identifié les emojis les plus communément utilisés pour chaque émotion¹⁰⁰. Le tableau A3 présente les résultats pour les 8 émotions utilisées dans cette étude.

98. P. Kralj Novak, J. Smailović, B. Sluban et I. Mozetič, « Sentiment of emojis », *PLoS one*, 7 décembre 2015.

99. http://kt.ijs.si/data/Emoji_sentiment_ranking/

100. A. Shoeb et G. De Melo, « Are emojis emotional ? A study to understand the association between emojis and emotions », *arXiv*, 2 mai 2020.

Tableau A2 – Emojis et sentiment
(tiré de P. Kraji Novak *et al*, 2015)

Emoji	Occurrences	Sentiment [- 1, + 1] ^a	Description
	14 622	0,221	Visage avec des larmes de joie
	15 194	0,746	Cœur
	6 359	0,678	Visage souriant avec des cœurs
	5 526	- 0,093	Visage pleurant bruyamment
	1 409	0,555	Biceps contracté
	756	- 0,173	Visage rouge de colère
	718	- 0,080	Visage endormi

a. - 1 correspondant à un sentiment négatif et 1, à un sentiment positif.

Tableau A3 – Emojis et émotions
(tiré de A. Shoeb et G. De Melo,2020)

Émotion	Sentiment associé	Emojis
Peur	0,97 1,00 0,90	 Visage apeuré  Visage criant de peur  Visage anxieux
Joie	1,00 0,94 0,94	 Visage souriant  Visage souriant avec bouche ouverte et yeux fermés  Visage avec des larmes de joie
Colère	1,00 1,00 0,75	 Visage colérique  Visage boudeur  Visage exaspéré
Tristesse	1,00 1,00 0,94	 Visage qui pleure  Visage pleurant abondamment  Cœur brisé
Dégoût	0,67 0,67 0,67	 Visage colérique  Visage déçu  Pouce vers le bas
Surprise	0,89 0,81 0,69	 Double point d'exclamation  Point d'exclamation  Visage figé par la peur
Anticipation	0,81 0,64 0,58	 Yeux  Nuage pensant  Sac d'argent
Confiance	0,83 0,75 0,72	 Visage qui envoie un baiser  Cœur  Rose

Annexe 2. La mesure du bien-être *via* Google et Facebook

En parallèle des travaux sur Twitter, des recherches ont été menées *via* les données disponibles sur Google Trends et Facebook¹⁰¹. Y. Algan *et al.* ont étudié les requêtes sur Google afin de mesurer l'évolution du bien-être¹⁰². Ils ont ainsi pu montrer l'intérêt d'utiliser les recherches effectuées sur Google en tant que complément aux mesures plus traditionnelles et cela, afin d'inférer l'évolution du bien-être subjectif. Plus récemment, A. Brodeur *et al.* ont utilisé l'évolution des recherches sur Google sur différents mots-clés (ennui, divorce, tristesse, stress, suicide, etc.) et ont identifié une augmentation du nombre des recherches sur ces mots-clés durant la pandémie de Covid-19¹⁰³. Dans cet article, les auteurs tentent d'identifier les effets du confinement sur le bien-être et la santé mentale en utilisant les recherches effectuées sur Google. Ainsi les pays ayant eu recours au confinement ont vu une forte augmentation de leurs recherches pour des mots tels qu'ennui (*boredom*), solitude (*loneliness*), inquiétude (*worry*), tristesse (*sadness*) et une diminution dans l'intensité des recherches pour des mots tels que stress, suicide ou divorce.

En ce qui concerne les données Facebook – qui contrairement aux données Google Trends ou Twitter ne sont pas publiques – les travaux de recherche ont été réalisés par des employés de Facebook à partir d'une base de données de messages anonymisés. A. Kramer a ainsi analysé les

101. Notons aussi les travaux de Reece et Danforth (2017) qui utilisent les photos publiées sur Instagram afin de prédire les risques de dépression chez les individus. Il n'existe pas d'indicateur agrégé du bien-être d'une population calculé à partir d'images ou de vidéos, mais la montée en puissance d'Instagram, de TikTok et des plateformes publiques permettant de partager du contenu (photo, vidéo) offre de nouvelles pistes aux chercheurs.

102. Art. cité.

103. A. Brodeur, D. Gray, A. Islam et S. Bhuiyan, « A literature review of the economics of COVID-19 », *Journal of Economic Surveys*, 35 (4), 2021.

messages Facebook (« *status updates* ») de plus de 100 millions d'utilisateurs aux États-Unis afin de créer un indicateur de bien-être¹⁰⁴. Kramer vérifie la validité de l'indicateur basé sur l'analyse de sentiment en comparant le bien-être calculé à partir des messages avec le bien-être déclaré par un échantillon de 1,324 utilisateurs ayant participé à un sondage. Il trouve ainsi une corrélation significative entre les deux variables, ce qui confirme la pertinence d'une mesure automatique *via* l'analyse textuelle. L'indicateur « Facebook Gross National Happiness » (FGNH) – calculé à partir de la différence quotidienne entre le nombre de *posts* positifs et le nombre de *posts* négatifs – a été utilisé ensuite par N. Wang *et al.*¹⁰⁵ qui ont comparé cet indicateur aux données récoltées par l'application Facebook de test de personnalité « MyPersonnalité »¹⁰⁶. Les auteurs ne trouvent quant à eux pas de corrélation entre l'indicateur FGNH et le bien-être subjectif collecté *via* leur application utilisée par plus de 20 000 utilisateurs.

104. A. Kramer, « An unobtrusive behavioral model of "gross national happiness" », in *Proceedings of the SIGCHI conference on human factors in computing systems*, 2010.

105. N. Wang, M. Kosinski, D. J. Stillwell et J. Rust, « Can well-being be measured using Facebook status updates ? Validation of Facebook's gross national happiness index », *Social Indicators Research*, 115 (1), 2014.

106. Pour la petite histoire, l'application My Personality – lancée par des chercheurs de l'Université de Cambridge – a été bannie en 2018 par Facebook à la suite du scandale de Cambridge Analytica, en raison de préoccupations concernant la manière dont les données collectées par l'application étaient utilisées.

Annexe 3. L'impact négatif causal de l'utilisation des réseaux sociaux sur le bien-être

L'utilisation d'un réseau social tel que Twitter n'est pas neutre et de nombreuses recherches se sont intéressées à l'impact de la fréquence d'utilisation des réseaux sociaux sur le bien-être individuel. À cet égard, deux théories peuvent s'opposer. Tout d'abord, les réseaux sociaux peuvent permettre de nouer des liens, de parler avec ses amis, de discuter. Des actions normalement associées à une hausse du bien-être. Mais les réseaux sociaux constituent aussi un lieu où se propagent les critiques, les moqueries, qui peuvent dégrader – par comparaison avec les autres – l'estime de soi. Les risques d'addiction et de déconnexion de la réalité posent aussi la question de l'impact sur la satisfaction dans la vie de l'utilisation des réseaux sociaux.

À travers l'utilisation d'un sondage en ligne, A. Masciantonio et *al.* cherchent à identifier les effets différenciés, pendant la période de la Covid-19, liés à l'utilisation des réseaux sociaux en distinguant à la fois le réseau social utilisé (Facebook, Instagram, Twitter et TikTok) et l'utilisation qui en est faite (active ou passive)¹⁰⁷. Ainsi, les auteurs trouvent que l'utilisation des réseaux sociaux a augmenté avec le confinement. De plus, à travers leurs résultats, il semblerait qu'une utilisation passive de Facebook tende à exercer un effet négatif sur le bien-être en raison d'un sentiment de comparaison avec les autres utilisateurs. Pour Twitter, deux effets existent. Le premier, lié à une utilisation active qui se traduit par un soutien social plus fort, ce qui tend à avoir un effet positif sur le bien-être (« *satisfaction with life* ») malgré un effet qui peut être également négatif du fait d'un soutien social (« *social support* ») négatif. Le deuxième

107. A. Masciantonio, D. Bourguignon, P. Bouchat, M. Balty, et B. Rimé, « Don't put all social network sites in one basket : Facebook, Instagram, Twitter, TikTok, and their relations with well-being during the Covid-19 pandemic », *PloS one*, 16 (3), 2021.

suggérant qu'une utilisation passive de Twitter semble aller de pair avec une réduction du bien-être en raison de la comparaison avec les autres utilisateurs.

Liste des figures, tableaux et encadré

Figures

Figure 1 – Nombre de tweets (échantillon de tweets en français).	21
Figure 2 – Indicateur Twitter du moral des Français (fréquence mensuelle)	27
Figure 3 – Indicateur Twitter du moral des Français (fréquence journalière	29
Figure 4 – Évolution du score de sentiment autour des chocs négatifs majeurs	30
Figure 5 – Évolution des indicateurs d'émotion (fréquence mensuelle). . .	32
Figure 6 – Score de sentiment moyen en fonction du jour de la semaine .	33
Figure 7 – Variation infra-journalière des émotions.	34
Figure 8 – Indicateur Twitter du moral des Français selon le genre	37
Figure 9 – Moral des Français sur Twitter par préférence politique	39
Figure 10 – Indice moyen de sentiment et médias.	42
Figure 11 – Indice moyen d'émotions et médias.	43
Figure 12 – Presse quotidienne nationale et émotions	44
Figure 13 – Hebdomadaires d'information et émotions	46
Figure 14 – Sélection de médias.	47
Figure 15 – Indicateur de moral des Français Twitter et Hedonometer. . .	49
Figure 16 – Évolution du nombre de tweets contenant l'emoji 📌 en 2020-2021	53
Figure 17 – Localisation des tweets de l'échantillon	55

Tableaux

Tableau 1 – Exemples de tweets classifiés avec l'approche VADER	24
Tableau 2 – Médias sur Twitter par type.	41

Tableau 3 – Hashtags les plus fréquemment associés à des émotions négatives (par an)	52
Tableau A 1 – Dictionnaire existant pour l'analyse de sentiment (tiré de A. Reagan et <i>al.</i> , 2017)	69
Tableau A 2 – Emojis et sentiment (tiré de P. Kraji Novak et <i>al.</i> , 2015).	77
Tableau A 3 – Emojis et émotions (tiré de A. Shoeb et G. De Melo, 2020)	78

Encadré

Encadré 1 – Amazon Mechanical Turk et la construction de dictionnaire de sentiment.	72
---	----

Bibliographie

- ABDAOUI, Amine, TCHECHMEDJIEV, Andon, DIGAN, William, BRINGAY, Sandra et JONQUET, Clément, « French ConText : Détecter la négation, la temporalité et le sujet dans les textes cliniques Français », *SIIM : Symposium sur l'ingénierie de l'information médicale*, 2017.
- ALGAN, Yann, BEASLEY, Elizabeth, COHEN, Daniel et FOUCAULT, Martial, *Les Origines du populisme, enquête sur un schisme politique et social*, Paris, Seuil et La République des Idées, 2019.
- ALLCOTT, Hunt et GENTZKOW, Matthew, « Social media and fake news in the 2016 election », *Journal of economic perspectives*, 31 (2), 2017.
- ALOTHALI, Eiman, ZAKI, Nazar, MOHAMED, Elfadil Abdalla et ALASHWAL, Hany, « Detecting social bots on Twitter : A literature review », in *2018 International conference on innovations in information technology (IIIT)*, 2018.
- BARBERÁ, Pablo, « Birds of the same feather tweet together : Bayesian ideal point estimation using Twitter data », *Political Analysis*, 2015.
- BELLET, Clément et FRIJTERS, Paul, « Big data and well-being », in *World Happiness Report*, 2019.
- CANOVA, Luciano, CURINI, Luigi et IACUS, Stefano, « Measuring idiosyncratic happiness through the analysis of Twitter : An application to the Italian case », *Social Indicators Research*, 121 (2), 2014.
- DAVIS, Clayton A., VAROL, Onur, FERRARA, Emilio, FLAMMINI, Alessandro et MENCZER, Filippo, « Botornot : A system to evaluate social bots », in *Proceedings of the 25th International Conference Companion on World Wide Web*, 2016.
- DÍAZ, Fernando et HENRÍQUEZ, Pablo A., « Social sentiment segregation : Evidence from Twitter and Google trends in Chile during the Covid-19 dynamic quarantine strategy », *PloS one*, 31 juillet 2021.
- DODDS, Peter S., HARRIS, Kameron H., BLISS, Catherine A. et DANFORTH, Christopher M., « Temporal patterns of happiness and information in a global social network : Hedonometrics and Twitter », *PloS one*, 7 décembre 2011.

- DURAHIM, Ahmet O. et COŞKUN, Mustafa, « #iamhappybecause : Gross national happiness through Twitter analysis and big data », *Technological Forecasting and Social Change*, 91 (9), octobre 2015.
- EINAV, Liran et LEVIN, Jonathan, « Economics in the age of big data », *Science*, 346 (6210), 7 novembre 2014.
- GIORGI, Salvatore, LYNN, Veronica, MATZ, Sandra, UNGAR, Lyle et SCHWARTZ, Andrew, « Correcting sociodemographic selection biases for accurate population prediction from social media », *arXiv*, 23 juillet 2021.
- GOLBECK, Jennifer et HANSEN, Derek, « A method for computing political preference among Twitter followers », *Social Networks*, 36, 2014.
- GOLDER, Scott, et MACY, Michael, « Diurnal and seasonal mood vary with work, sleep, and daylength across diverse cultures », *Science*, 333 (6051), 30 septembre 2011.
- HALLDORSDDOTTIR, Thorhildur, THORISDDOTTIR, Ingibjorg, MEYERS, Caine, ASGEIRSDOTTIR, Bryndis, KRISTJANSSON, Alfgeir, VALDIMARSDOTTIR, Heiddis et SIGFUSDDOTTIR, Inga, « Adolescent well-being amid the Covid-19 pandemic : Are girls struggling more than boys ? », *JCPP advances*, 1 (2), juillet 2021.
- HAO, Bibo, LI, Lin, GAO, Rui, LI, Ang et ZHU, Tingshao, « Sensing subjective well-being from social media », in *Proceedings of Active Media Technology, 10th International Conference, AMT 2014*, Warsaw, Poland, August 11-14, 2014.
- HUTTO, C. et GILBERT, Éric, « VADER : A parsimonious rule-based model for sentiment analysis of social media text », *Proceedings of the International AAAI*, 2014.
- IACUS, Stefano et PORRO, Giuseppe, *Subjective Well-Being and Social Media*, Londres, Chapman and Hall/CRC, 2021.
- IACUS, Stefano, PORRO, Giuseppe, SALINI, Silvia et SILETTI, Elena, « Controlling for selection bias in social media indicators through official statistics : A proposal », *Journal of Official Statistics*, 36 (2), 2020.
- IACUS, Stefano, PORRO, Giuseppe, SALINI, Silvia et SILETTI, Elena, « An Italian composite subjective well-being index : The voice of Twitter users from 2012 to 2017 », *Social Indicators Research*, 161 (2), juin 2022.

- JAIDKA, Kokil, GIORGI, Salvatore, SCHWARTZ, H. Andrew, KERN, Margaret, UNGAR, Lyle et EICHSTAEDT, Johannes, « Estimating geographic subjective well-being from Twitter : A comparison of dictionary and data-driven language methods », *Proceedings of the National Academy of Sciences*, 2020.
- KOSINSKI, Michal, STILLWELL, David et GRAEPEL, Thore, « Private traits and attributes are predictable from digital records of human behavior », *Proceedings of the National Academy of Sciences*, 2013.
- LAZER, David, KENNEDY, Rian, KING, Gary et VESPIGNANI, Alessandro, « The parable of Google flu : traps in big data analysis », *Science*, 343 (6176), 14 mars 2014.
- LUDU, Puneet, « Inferring gender of a twitter user using celebrities it follows », *arXiv preprint*, 2014.
- LWIN, May Oo, LU, Jiahui, SHELDENKAR, Anita, SCHULZ, Peter J., SHIN, Wonsun, GUPTA, Raj et YANG, Yinping, « Global sentiments surrounding the Covid-19 pandemic on Twitter : Analysis of Twitter trends », *JMIR public health and surveillance*, 6 (2), 17 avril 2020.
- METZEL, Hannah, PELLERT, Max et GARCIA, David, « Using social media data to capture emotions before and during Covid-19 », *World Happiness Report*, 2022.
- MISLOVE, Alan, LEHMANN, Sule, AHN, Yong-Yeol, ONNELA, Jukka Pekka et ROSENQUIST, Niels, « Understanding the demographics of Twitter users », in *Proceedings of the International AAAI Conference on Web and Social Media*, 2011.
- MITCHELL, Lewis, FRANK, Morgan, HARRIS, Kameron, DODDS, Peter et DANFORTH, Christopher, « The geography of happiness : Connecting twitter sentiment and expression, demographics, and objective characteristics of place », *PloS one*, 2013.
- NIENHUIS, Carl et LESSER, Iris, « The impact of Covid-19 on women's physical activity behavior and mental well-being », *International Journal of Environmental Research and Public Health*, 17 (23), 4 décembre 2020.
- QUERCIA, Daniele, ELLIS, Jonathan, CAPRA, Licia et CROWCROFT, Jon, « Tracking gross community happiness from tweets », in *Proceedings of the ACM 2012 Conference on Computer Supported Cooperative Work*, 2012.

- REAGAN, Andrew J., TIVNAN, Brian, WILLIAMS, Jake Ryland, DANFORTH, Christopher et DODDS, Peter, « Benchmarking sentiment analysis methods for large-scale texts : A case for using continuum-scored words and word shift graphs », *arXiv*, 2016.
- SAP, Maarten, PARK, Gregory, EICHSTAEDT, Johannes, KERN, Margaret, STILLWELL, David, KOSINSKI, Michal, UNGAR, Lyle et SCHWARTZ, Hansen, « Developing age and gender predictive lexica over social media », in *Proceedings of the 2014 Conference on Empirical Methods in Natural Language Processing*, 2014.
- SARRACINO, Francesco, GREYLING, Talita, O'CONNOR, Kelsey, PERONI, Chiara et ROSSOUW, Stephanie, « Trust predicts compliance to Covid-19 containment policies : Evidence from ten countries using big data », Department of Economics, University of Siena, 2021.
- SCHWARTZ, Hansen, EICHSTAEDT, Johannes, KERN, Margaret, DZIURZYNSKI, Lukasz, LUCAS, Richard, AGRAWAL, Megha et UNGAR, Lyle, « Characterizing geographic variation in well-being using tweets », in *Seventh International AAAI Conference on Weblogs and Social Media*, 2013.
- SCHWARTZ, Hansen, EICHSTAEDT, Johannes, KERN, Margaret, DZIURZYNSKI, Lukasz, RAMONES, Stephanie, AGRAWAL, Megha et UNGAR, Lyle, « Personality, gender, and age in the language of social media : The open-vocabulary approach », *PLoS one*, 2013.
- SHOEB, Abu Awal et DE MELO, Gerard, « EmoTag1200 : Understanding the association between emojis and emotions », *Association for Computational Linguistics*, 2020.
- STIEGLITZ, Stefan, BRACHTEN, Florian, BERTHELÉ, Davina, SCHLAUS, Mira, VENETOPOULOU, Chrissoula et VEUTGEN, Daniel, « Do social bots (still) act different to humans ? : Comparing metrics of social bots with those of humans », in *International Conference on Social Computing and Social Media*, Springer, 2017.
- VOUKELATOU, Vassiliki, GABRIELLI, Lorenzo, MILIOU, Ioanna, CRESCI, Stefano, SHARMA, Rajesh, TESCONI, Maurizio et PAPPALARDO, Luca, « Measuring objective and subjective well-being : Dimensions and data sources », *International Journal of Data Science and Analytics*, 11, 29 juin 2020.

- WALDHERR, Annie, MAIER, Daniel, MILTNER, Peter et GÜNTHER, Enrico, « Big data, big noise : The challenge of finding issue networks on the web », *Social Science Computer Review*, 35 (4), août 2017.
- WANG, Yilin., WU, Peijing, LIU, Xiaoqian., LI, Sijia, ZHU, Tingshao et ZHAO Nan, « Subjective well-being of Chinese Sina Weibo users in residential lockdown during the Covid-19 pandemic : Machine learning analysis », *Journal of Medical Internet Research*, décembre 2020.
- XUE, Jia, CHEN, Junxiang, HU, Ran, CHEN, Chen, ZHENG, Chengda, SU, Yue et ZHU, Tingshao, « Twitter discussions and emotions about the Covid-19 pandemic : Machine learning approach », *Journal of Medical Internet Research*, 22 (11), novembre 2020.
- YANG, Chao et SRINIVASAN, Padmini, « Life satisfaction and the pursuit of happiness on Twitter », *PLoS one*, 2016.
- YANG, Kai- Cheng, FERRARA, Emilio et MENCZER, Filippo, « Botometer 101 : Social bot practicum for computational social scientists », arXiv, 5 janvier 2022.

ORGANIGRAMME DU CEPREMAP

Président : Benoît Cœuré
Directeur : Daniel Cohen
Directrice adjointe : Claudia Senik

MACROÉCONOMIE

Observatoire macroéconomie

François Langot

Gilles Saint-Paul

Thomas Brand

(directeur exécutif)

Projet Dynare

Stéphane Adjémian

Sébastien Villemot

Projet DbNomics

Thomas Brand

BIEN-ÊTRE, EMPLOI ET POLITIQUES PUBLIQUES

Observatoire du bien-être

Yann Algan

Andrew Clark

Sarah Flèche

Claudia Senik (directrice)

Mathieu Perona (directeur exécutif)

Travail et emploi

Luc Behaghel

Philippe Askenazy

Dominique Meurs

Économie publique et redistribution

Maya Bacache-Beauvallet

Antoine Bozio

Brigitte Dormont

MONDIALISATION, DÉVELOPPEMENT ET ENVIRONNEMENT

Observatoire mondialisation

Miren Lafourcade

Sylvie Lambert

Katheline Schubert

Ishac Diwan

(directeur exécutif)

Groupe Chine-Inde

Guilhem Cassan

Maelys de la Rupelle

Clément Imbert

Oliver Vanden Eynde

Thomas Vendryes

Méditerranée-Moyen-Orient

Ishac Diwan

DANS LA MÊME COLLECTION

La Lancinante Réforme de l'assurance maladie, par Pierre-Yves Geoffard, 2006, 48 pages.

La Flexicurité danoise. Quels enseignements pour la France ?, par Robert Boyer, 2007, 3^e tirage, 54 pages.

La Mondialisation est-elle un facteur de paix ?, par Philippe Martin, Thierry Mayer et Mathias Thoenig, 2006, 2^e tirage, 56 pages.

L'Afrique des inégalités : où conduit l'histoire, par Denis Cogneau, 2007, 64 pages.

Électricité : faut-il désespérer du marché ?, par David Spector, 2007, 2^e tirage, 56 pages.

Une jeunesse difficile. Portrait économique et social de la jeunesse française, par Daniel Cohen (éd.), 2007, 238 pages.

Les Soldes de la loi Raffarin. Le contrôle du grand commerce alimentaire, par Philippe Askenazy et Katia Weidenfeld, 2007, 60 pages.

La Réforme du système des retraites : à qui les sacrifices ?, par Jean-Pierre Laffargue, 2007, 52 pages.

Les Pôles de compétitivité. Que peut-on en attendre ?, par Gilles Duranton, Philippe Martin, Thierry Mayer et Florian Mayneris, 2008, 2^e tirage, 84 pages.

Le Travail des enfants. Quelles politiques pour quels résultats ?, par Christelle Dumas et Sylvie Lambert, 2008, 82 pages.

Pour une retraite choisie. L'emploi des seniors, par Jean-Olivier Hairault, François Langot et Theptida Sopraseuth, 2008, 72 pages.

La Loi Galland sur les relations commerciales. Jusqu'où la réformer ?, par Marie-Laure Allain, Claire Chambolle et Thibaud Vergé, 2008, 74 pages.

Pour un nouveau système de retraite. Des comptes individuels de cotisations financés par répartition, par Antoine Bozio et Thomas Piketty, 2008, 2^e tirage, 100 pages.

Les Dépenses de santé. Une augmentation salubre ?, par Brigitte Dormont, 80 pages, 2009.

De l'euphorie à la panique. Penser la crise financière, par André Orléan, 2009, 3^e tirage, 112 pages.

Bas salaires et qualité de l'emploi : l'exception française ?, par Ève Caroli et Jérôme Gautié (éd.), 2009, 510 pages.

Pour la taxe carbone. La politique économique face à la menace climatique, par Katheline Schubert, 2009, 92 pages.

Le Prix unique du livre à l'heure du numérique, par Mathieu Perona et Jérôme Pouyet, 2010, 92 pages.

Pour une politique climatique globale. Blocages et ouvertures, par Roger Guesnerie, 2010, 96 pages.

Comment faut-il payer les patrons ?, par Frédéric Palomino, 2011, 74 pages.

Portrait des musiciens à l'heure du numérique, par Maya Bacache-Beauvallet, Marc Bourreau et François Moreau, 2011, 94 pages.

L'Épargnant dans un monde en crise. Ce qui a changé, par Luc Arrondel et André Masson, 2011, 112 pages.

Handicap et dépendance. Drames humains, enjeux politiques, par Florence Weber, 2011, 76 pages.

Les Banques centrales dans la tempête. Pour un nouveau mandat de stabilité financière, par Xavier Ragot, 2012, 80 pages.

L'Économie politique du néolibéralisme. Le cas de la France et de l'Italie, par Bruno Amable, Elvire Guillaud et Stefano Palombarini, 2012, 164 pages.

Faut-il abolir le cumul des mandats ?, par Laurent Bach, 2012, 126 pages.

Pour l'emploi des seniors. Assurance chômage et licenciements, par Jean-Olivier Hairault, 2012, 78 pages.

L'État-providence en Europe. Performance et dumping social, par Mathieu Lefebvre et Pierre Pestieau, 80 pages, 2012.

Obésité. Santé publique et populisme alimentaire, par Fabrice Étilé, 2013, 124 pages.

La Discrimination à l'embauche sur le marché du travail français, par Nicolas Jacquemet et Anthony Edo, 2013, 78 pages.

Travailler pour être aidé ? L'emploi garanti en Inde, par Clément Imbert, 2013, 74 pages.

Hommes/Femmes. Une impossible égalité professionnelle ?, par Dominique Meurs, 2014, 106 pages.

Le Fédéralisme en Russie ? Les leçons de l'expérience internationale, par Ekaterina Zhuravskaya, 2014, 68 pages.

Bien ou mal payés ? Les travailleurs du public et du privé jugent leurs salaires, par Christian Baudelot, Damien Cartron, Jérôme Gautié, Olivier Godechot, Michel Gollac et Claudia Senik, 2014, 232 pages.

La Caste dans l'Inde en développement. Entre tradition et modernité, par Guilhem Cassan, 2015, 72 pages.

Libéralisation, innovation et croissance. Faut-il les associer ?, par Bruno Amable et Ivan Ledezma, 2015, 122 pages.

Les Allocations logement. Comment les réformer ?, par Antoine Bozio, Gabrielle Fack et Julien Grenet (dir.), 2015, 98 pages.

Avoir un enfant plus tard. Enjeux sociodémographiques du report des naissances, par Hippolyte d'Albis, Angela Greulich et Grégory Ponthière, 2015, 128 pages.

La Société de défrance. Comment le modèle social français s'autodétruit, par Yann Algan et Pierre Cahuc, 2016, 2^e édition, 110 pages.

Leçons de l'expérience japonaise. Vers une autre politique économique ?, par Sébastien Lechevalier et Brieuc Monfort, 2016, 228 pages.

Filles + sciences = une équation insoluble ? Enquêtes sur les classes préparatoires scientifiques, par Marianne Blanchard, Sophie Orange et Arnaud Pierrel, 2016, 152 pages.

Qualité de l'emploi et productivité, par Philippe Askenazy et Christine Erhel, 2017, 104 pages.

En finir avec les ghettos urbains ? Retour sur l'expérience des zones franches urbaines, par Miren Lafourcade et Florian Mayneris, 2017, 136 pages.

Repenser l'immigration en France, par Hillel Rapoport, 2018, 102 pages.

Les Français, le bonheur et l'argent, par Yann Algan, Elizabeth Beasley et Claudia Senik, 2018, 80 pages.

La Transition écologique en Chine. Mirage ou « virage vert » ?, par Stéphanie Monjon et Sandra Poncet, 2018, 176 pages.

Biens publics, charité privée. Comment l'État peut-il réguler le charity business ?, par Gabrielle Fack, Camille Landais et Alix Myczkowski, 2018, 118 pages.

Competition between hospitals. Does it affect quality of care ?, par Brigitte Dormont et Carine Milcent (éd.), 2018, 236 pages.

La Polarisation de l'emploi en France. Ce qui s'est aggravé depuis la crise de 2008, par Ariell Reshef et Farid Toubal, 2019, 96 pages.

Voter autrement, par Jean-François Laslier, 2019, 140 pages.

Mondialisation des échanges et protection des consommateurs. Comment les concilier ?, par Anne-Célia Disdier, 2020, 108 pages.

Comment lutter contre la fraude fiscale ? Les enseignements de l'économie comportementale, par Nicolas Jacquemet, Stéphane Luchini et Antoine Malézieux, 2020, 104 pages.

Remédier aux déserts médicaux, par Magali Dumontet et Guillaume Chevillard, 2020, 126 pages.

Comment les garçons ? L'économie du football féminin, par Luc Arrondel et Richard Duhautois, préface d'Hervé Mathoux, 2020, 184 pages.

La Valeur des réseaux. Économie des interactions sociales, par Margherita Comola, 2020, 76 pages.

La Transition énergétique : objectif ZEN, par Fanny Henriot et Katheline Schubert, 2021, 122 pages.

L'Hélicoptère monétaire. Au-delà du mythe, par Florin Bilbiie, Alais Martin-Baillon et Gilles Saint-Paul, 2021, 110 pages.

Mise en pages
TyPAO sarl
75011 Paris

Imprimerie Maury
N° d'impression : *****
Dépôt légal : juillet 2022