

JOURNAL
DE
MATHÉMATIQUES

PURES ET APPLIQUÉES

FONDÉ EN 1836 ET PUBLIÉ JUSQU'EN 1874

PAR JOSEPH LIOUVILLE

H. POINCARÉ

Sur les courbes définies par les équations différentielles

Journal de mathématiques pures et appliquées 4^e série, tome 1 (1885), p. 167-244.

http://www.numdam.org/item?id=JMPA_1885_4_1__167_0



NUMDAM

Article numérisé dans le cadre du programme
Gallica de la Bibliothèque nationale de France
<http://gallica.bnf.fr/>

et catalogué par Mathdoc
dans le cadre du pôle associé BnF/Mathdoc
<http://www.numdam.org/journals/JMPA>

Sur les courbes définies par les équations différentielles

(TROISIÈME PARTIE) :

PAR M. H. POINCARÉ.

Les deux premières Parties de ce travail ont paru dans ce Journal, la première aux mois de novembre et décembre 1881 (3^e série, t. VII) et la deuxième au mois d'août 1882 (3^e série, t. VIII). Dans ce qui va suivre, je conserverai, malgré les objections auxquelles elles peuvent donner lieu, les dénominations employées dans les deux premières Parties, afin de n'avoir pas à donner de définitions nouvelles.

CHAPITRE X.

STABILITÉ ET INSTABILITÉ.

On n'a pu lire les deux premières Parties de ce Mémoire sans être frappé de la ressemblance que présentent les diverses questions qui y sont traitées avec le grand problème astronomique de la stabilité du système solaire. Ce dernier problème est, bien entendu, beaucoup plus compliqué, puisque les équations différentielles du mouvement des corps célestes sont d'ordre très élevé. Il y a même plus, on rencontrera, dans ce problème, une difficulté nouvelle, essentiellement différente de celles que nous avons eu à surmonter dans l'étude du premier ordre, et j'ai l'intention de la faire ressortir, sinon dans cette troisième Partie, du moins dans la suite de ce travail.

Mais, quoi qu'il en soit, si la solution exige de plus grands efforts et des procédés nouveaux, l'analogie des questions à résoudre n'en est pas moins évidente. Pour étudier l'équation différentielle

$$\frac{dx}{X} = \frac{dy}{Y},$$

on peut poser

$$\frac{dx}{dt} = X, \quad \frac{dy}{dt} = Y.$$

Regardant ensuite x et y comme les coordonnées d'un point mobile, t comme le temps, on a à rechercher quel est le mouvement d'un point dont on donne la vitesse en fonction de ses coordonnées. C'est ce mouvement que nous avons étudié, et nous avons cherché à résoudre des questions, telles que celles-ci; le point mobile décrira-t-il une courbe fermée? Restera-t-il toujours à l'intérieur d'une certaine portion du plan? En d'autres termes, et pour parler le langage astronomique, nous avons recherché si l'orbite de ce point était stable ou instable.

Nous pourrions nous poser des questions analogues, lorsque X et Y ne seront plus des polynômes en x et y , mais des fonctions algébriques de ces variables, ou bien encore lorsque l'équation différentielle sera d'ordre supérieur au premier.

Mais auparavant il importe de définir exactement ce qu'on doit entendre par stabilité ou instabilité. Pour cela, nous allons étudier les cinq équations qui suivent et qui nous donneront des exemples de tous les cas qui peuvent se présenter :

1^o Soient d'abord

$$\frac{dx}{dt} = -y, \quad \frac{dy}{dt} = x,$$

il vient, pour l'équation de la trajectoire,

$$x^2 + y^2 = \text{const.}$$

Les trajectoires étant des courbes fermées, la stabilité est complète.

2^o Soient maintenant, en coordonnées polaires,

$$\frac{d\omega}{dt} = h, \quad \frac{dz}{dt} = \sqrt{1 - (z - 2)^2},$$

h étant une constante quelconque. Il serait aisé, d'ailleurs, de passer de cette équation en coordonnées polaires à l'équation correspondante en coordonnées rectangulaires, et l'on verrait que, si l'on met cette équation sous la forme

$$\frac{dx}{dt} = X, \quad \frac{dy}{dt} = Y,$$

X et Y ne seront plus des polynômes entiers en x et y , mais des fonctions algébriques de ces variables. L'équation différentielle sera encore du premier ordre, mais sera de degré supérieur.

On trouve immédiatement l'intégrale

$$\varphi = 2 + \sin\left(\frac{\omega}{h} + C\right),$$

C étant une constante d'intégration.

Si h est commensurable avec 2π , la trajectoire est une courbe fermée, et l'on retombe sur le cas précédent. Supposons qu'il n'en soit pas ainsi.

Il arrivera alors que la trajectoire ne sera pas une courbe fermée; mais néanmoins elle jouira d'une certaine stabilité : on peut même dire d'une certaine périodicité d'une nature particulière. En effet, soit M un point de la trajectoire, occupé un temps t par le point mobile. Décrivons autour du point M un cercle de rayon r aussi petit que nous voudrions. Le point mobile partant du point M sortira évidemment de ce cercle, mais il viendra traverser de nouveau ce petit cercle une *infinité de fois*, et cela, quelque petit que soit r . En d'autres termes, le point mobile partant du point M ne pourra jamais revenir en ce point, mais il reviendra en des points infiniment voisins de M .

En second lieu, le point M restera toujours à l'intérieur de la couronne limitée par les deux cercles

$$\varphi = 1 \quad \text{et} \quad \varphi = 3.$$

Mais sa trajectoire remplira entièrement cette couronne, sans laisser de lacune. Je veux dire que, dans toute aire plane, si petite qu'elle soit, située à l'intérieur de la couronne, il y a des points de la trajectoire.

Les Allemands diraient que la *Punktmenge*, formée par les différents points de la trajectoire, est *überalldicht* à l'intérieur de la couronne.

Ce second cas ne peut pas se présenter pour les équations différentielles du premier ordre et du premier degré. C'est pourquoi nous ne l'avons jamais rencontré jusqu'ici

2° Soient maintenant, en coordonnées polaires,

$$\frac{d\rho}{dt} = 1 + \rho^2, \quad \frac{d\omega}{dt} = \frac{1}{h}$$

ou, en coordonnées rectangulaires,

$$\frac{dx}{dt} = \frac{1+x^2+y^2}{\sqrt{x^2+y^2}}x - \frac{y}{h}, \quad \frac{dy}{dt} = \frac{1+x^2+y^2}{\sqrt{x^2+y^2}}y + \frac{x}{h}.$$

L'intégrale générale est

$$\rho = \tan(h\omega + C),$$

C étant une constante d'intégration.

Ici encore, si h est incommensurable avec 2π , le point mobile ne peut jamais revenir à son point de départ, mais il peut revenir en des points infiniment voisins.

La différence avec le cas précédent, c'est que le point mobile n'est plus assujéti à rester dans une certaine région du plan, et que c'est le plan tout entier que la trajectoire remplit sans lacune (*überalldicht*).

Ce troisième cas ne peut, pas plus que le deuxième, se présenter pour les équations du premier ordre et du premier degré

4° Comme quatrième exemple, nous prendrons la spirale logarithmique dont l'équation différentielle s'écrit

$$\frac{d\rho}{d\omega} = m\rho$$

ou bien

$$\frac{dx}{d\omega} = mx - y, \quad \frac{dy}{d\omega} = my + x.$$

L'intégrale générale est, comme on sait,

$$\rho = Ce^{m\omega}.$$

Soit M le point de départ du point mobile; si nous décrivons autour du point M un cercle de rayon suffisamment petit, nous verrons le point mobile, partant de M, sortir de ce cercle, et, après en être sorti, *n'y plus jamais rentrer*. C'est le contraire de ce qui se passait dans les trois cas examinés plus haut, et où le point mobile, après être sorti d'un cercle très petit, y rentrait ensuite une infinité de fois. A ce point de vue, on peut dire que la trajectoire est instable.

On sait que les cercles $\rho = \text{const.}$ sont, pour nos trajectoires, des *cycles sans contact*. De plus, tout le plan est sillonné par ces cycles sans contact, sans qu'il y ait de *cycles limites*. C'est à la présence de ces cycles sans contact qu'est due l'instabilité de la trajectoire.

5° Soit enfin l'équation

$$\frac{dz}{d\omega} = (\rho - 1)(\rho - 2).$$

Les cercles $\rho = \text{const.}$ sont encore des cycles sans contact, excepté les cercles $\rho = 1$ et $\rho = 2$, qui sont des cycles limites. L'intégrale générale étant

$$\rho = \frac{Ce^{\omega} - 2}{Ce^{\omega} - 1},$$

on voit aisément que la trajectoire est instable, c'est-à-dire que, après être sortie d'un cercle suffisamment petit décrit autour du point de départ, elle ne pourra plus y rentrer.

La différence avec le cas précédent tient à l'existence des cycles limites. Il en résulte que, si le point de départ est à l'intérieur de la couronne limitée par les deux cercles $\rho = 1$ et $\rho = 2$, le point mobile restera toujours à l'intérieur de cette couronne.

Nous pouvons maintenant, en nous référant aux exemples précédents, donner une définition précise de la stabilité. Nous dirons que la trajectoire d'un point mobile est stable, lorsque, décrivant autour du point de départ un cercle ou une sphère de rayon r , le point mobile, après être sorti de ce cercle ou de cette sphère, y rentrera une infinité de fois, et cela, quelque petit que soit r . C'est ce qui arrive dans les trois premiers exemples.

Elle sera instable si, après être sorti de ce cercle ou de cette sphère,

le point mobile n'y rentre plus. C'est ce qui arrive dans les deux derniers exemples.

La stabilité ainsi définie n'a qu'une importance théorique. Pour la pratique, il faudrait déterminer une région de l'espace où le point mobile reste constamment renfermé. Il arrive justement que la détermination d'une pareille région est beaucoup plus difficile dans le cas de la stabilité que dans le cas de l'instabilité. Il y a là une difficulté, sur laquelle je ne veux pas insister dans ce moment, mais qui fera, dans la suite de ce travail, l'objet d'assez longs développements.

Dans les cas que nous avons étudiés jusqu'ici, c'est-à-dire pour les équations du premier ordre et du premier degré, les trajectoires sont des cycles, c'est-à-dire des courbes fermées ou des spirales (*voir* théorème XII, t. VIII, p. 255). Dans le premier cas, elles sont stables; dans le second, instables. On ne peut donc jamais rencontrer rien de semblable à ce que nous venons d'observer dans les deuxième et troisième exemples.

En général, le plan est sillonné d'une infinité de cycles sans contact et de cycles limites (théorème XVIII, t. VIII, p. 274). Toutes les trajectoires sont des spirales, excepté les cycles limites.

L'instabilité est donc la règle, et la stabilité l'exception.

Il peut arriver aussi, dans des cas très exceptionnels, que le plan, au lieu d'être sillonné par une infinité de cycles sans contact, est sillonné par une infinité de courbes fermées satisfaisant à l'équation différentielle proposée et formant un de ces systèmes que nous avons appelés *topographiques* (t. VII, p. 383). Il y a alors stabilité. C'est ce qui arrive en particulier dans le voisinage de ces points singuliers exceptionnels que j'ai appelés *centres* (t. VII, p. 391) et dont nous allons faire une étude plus approfondie.

CHAPITRE XI.

THÉORIE DES CENTRES.

Écrivons notre équation différentielle sous la forme

$$\frac{dx}{dt} = X, \quad \frac{dy}{dt} = Y,$$

de manière à lui faire représenter le mouvement d'un point mobile. Supposons que X et Y sont des polynômes de degré n en x et en y . Supposons de plus que nous ayons pris pour origine le point singulier que nous voulons étudier, de telle façon que X et Y s'annulent avec x et y . Nous pourrions écrire alors

$$\begin{aligned} X &= X_1 + X_2 + \dots + X_n, \\ Y &= Y_1 + Y_2 + \dots + Y_n, \end{aligned}$$

X_i et Y_i étant des polynômes homogènes de degré i en x et en y . Cherchons maintenant si l'on peut former un polynôme F en x et en y , de telle façon que, dans l'expression

$$\Phi = \frac{dF}{dx} X + \frac{dF}{dy} Y,$$

qui est aussi un polynôme entier en x et en y , les termes de degré inférieur à p en x et en y soient tous nuls.

Nous pourrions écrire

$$F = F_2 + F_3 + F_4 + \dots$$

(F_i étant homogène de degré i en x et en y), en supposant, ce qui est nécessaire pour notre objet, que F ne contienne pas de terme de degré 0 ou 1.

Posons ensuite

$$\Phi_{ik} = \frac{dF_i}{dx} X_k + \frac{dF_i}{dy} Y_k;$$

Φ_{ik} sera un polynôme homogène de degré $i + k - 1$.

Si nous écrivons que, dans Φ , tous les termes de degré inférieur à p sont nuls, il viendra

$$(1) \quad \left\{ \begin{array}{l} \Phi_{21} = 0, \\ \Phi_{31} = -\Phi_{22}, \\ \Phi_{41} = -\Phi_{32} - \Phi_{23}, \\ \Phi_{51} = -\Phi_{42} - \Phi_{33} - \Phi_{24}, \\ \dots\dots\dots \\ \Phi_{p-11} = -\Phi_{p-22} - \Phi_{p-33} - \dots - \Phi_{2p-2}. \end{array} \right.$$

La première de ces équations nous donnera F_2 , la deuxième F_3 , ..., et enfin la $p - 2^{\text{ième}}$ nous donnera F_{p-1} , *pourvu toutefois qu'il soit possible d'y satisfaire.*

Considérons d'abord la première équation.

Soient

$$F_2 = ax^2 + 2bxy + cy^2,$$

$$X_1 = \alpha x + \beta y, \quad Y_1 = \gamma x + \delta y,$$

il vient

$$\frac{1}{2}\Phi_{2,1} = (ax + by)(\alpha x + \beta y) + (bx + cy)(\gamma x + \delta y),$$

de sorte que la première équation (1) entraîne les trois suivantes :

$$2) \quad \begin{cases} a\alpha + b\gamma = 0, \\ a\beta + b(\alpha + \delta) + c\gamma = 0, \\ b\beta + c\delta = 0; \end{cases}$$

ces équations ne sont compatibles que si l'on a

$$3) \quad \begin{vmatrix} \alpha & 0 & \gamma \\ \beta & \alpha + \delta & \gamma \\ 0 & \beta & \delta \end{vmatrix} = 0.$$

Nous supposons de plus que les valeurs de a , b et c , que l'on tire des équations (2), sont telles que la forme F_2 soit définie positive.

Si ces deux conditions sont remplies, on retombera sur le *quatrième cas subordonné* dont nous avons dit quelques mots à la page 390 du Tome VII, mais que nous n'avons pas encore étudié à fond.

Si elles n'étaient pas remplies, au contraire, le point singulier serait un nœud, un foyer ou un col, et nous n'aurions rien à ajouter à ce que nous avons déjà dit au sujet de ces points. Nous supposons donc ces deux conditions satisfaites.

La forme F_2 étant définie positive, nous pouvons toujours l'écrire sous la forme suivante :

$$(\lambda x + \mu y)^2 + (\lambda' x + \mu' y)^2.$$

Si nous changeons ensuite de variables en posant

$$x' = (\lambda x + \mu y), \quad y' = (\lambda' x + \mu' y),$$

il viendra

$$F_2 = x'^2 + y'^2.$$

En conséquence, nous pouvons toujours supposer que

$$F_2 = x^2 + y^2;$$

car, si cela n'était pas, il suffirait d'un changement linéaire de variables pour ramener F_2 à cette forme. Nous ferons désormais cette hypothèse.

Quelles en sont les conséquences au sujet de X_1 et de Y_1 ?

Les équations (2) deviennent

$$\alpha = 0, \quad \delta = 0, \quad \beta + \gamma = 0;$$

d'où

$$X_1 = \beta y, \quad Y_1 = -\beta x.$$

Les autres équations (1) s'écrivent alors

$$(4) \quad y \frac{dF_q}{dx} - x \frac{dF_q}{dy} = H_q,$$

où F_q est un polynôme homogène de degré q qu'il s'agit de déterminer, pendant que H_q est un polynôme homogène de degré q qu'on peut considérer comme donné, puisqu'il ne dépend que des polynômes X et Y et des polynômes homogènes $F_2, F_3, \dots, F_{q-1}$ que l'on a dû calculer avant F_q .

Dans quel cas est-il possible de satisfaire à une équation de la forme (4)?

Pour résoudre cette question, nous allons passer aux coordonnées polaires en posant

$$x = \rho \cos \omega, \quad y = \rho \sin \omega.$$

Il viendra

$$y \frac{dF}{dx} - x \frac{dF}{dy} = - \frac{dF}{d\omega}$$

et

$$F_q = \rho^q \varphi(\omega), \quad H_q = \rho^q \psi(\omega).$$

De plus, on aura

$$\varphi(\omega) = \sum A_k \cos k\omega + \sum B_k \sin k\omega,$$

$$\psi(\omega) = \sum C_k \cos k\omega + \sum D_k \sin k\omega.$$

Dans ces expressions, k ne pourra prendre que des valeurs inférieures ou au plus égales à q , et de même parité que q . En particulier, si q est impair, il ne pourra pas y avoir de terme tout connu (c'est-à-dire de terme où $k = 0$).

L'équation (4) s'écrit alors

$$-\frac{dz}{d\omega} = \psi(\omega).$$

Pour qu'on puisse y satisfaire, il faut et il suffit que $\psi(\omega)$ ne contienne pas de terme tout connu, c'est-à-dire que l'on ait

$$C_0 = 0.$$

Cette condition est remplie d'elle-même, comme nous venons de le voir, lorsque q est impair. Elle ne l'est pas, au contraire, en général, lorsque q est pair.

Si q est impair, il y a donc toujours une manière et une seule de satisfaire à l'équation (4); il suffit de poser

$$A_k = \frac{D_k}{k}, \quad B_k = -\frac{C_k}{k}.$$

Si q est pair et si, cependant, C_0 est nul, il y a une infinité de manières de satisfaire à notre équation; on fera encore, en effet,

$$A_k = \frac{D_k}{k}, \quad B_k = -\frac{C_k}{k}$$

si k est différent de zéro, et l'on pourra choisir A_0 arbitrairement.

Qu'arrive-t-il enfin si q est pair et si C_0 n'est pas nul? Dans ce cas, il est impossible de satisfaire à l'équation (4), mais on peut choisir F_q , de telle façon que

$$y \frac{dF_q}{dx} - x \frac{dF_q}{dy} < H_q,$$

pour toutes les valeurs de x et de y si C_0 est positif.

Au contraire, si C_0 est négatif, on pourra choisir F_q , de telle façon que l'on ait toujours

$$y \frac{dF_q}{dx} - x \frac{dF_q}{dy} > H_q.$$

Il suffit, pour cela, de faire

$$A_k = \frac{D_k}{k}, \quad B_k = -\frac{C_k}{k}$$

pour toutes les valeurs de k différentes de zéro, A_0 restant arbitraire. Il vient alors

$$\psi(\omega) + \frac{dz}{d\omega} = C_0$$

ou bien

$$H_q - y \frac{dF_q}{dx} + x \frac{dF_q}{dy} = C_0(x^2 + y^2)^{\frac{q}{2}}.$$

Soient, par exemple,

$$q = 4, \quad H_q = \beta x^4 - 2\alpha x^2 y^2 + \beta y^4.$$

Il viendra

$$H_4 = z^4 \left(\frac{3\beta - \alpha}{4} + \frac{\alpha + \beta}{4} \cos 4\omega \right).$$

Faisons alors

$$F_4 = z^4 \left(-\frac{\alpha + \beta}{16} \sin 4\omega \right) = -\frac{\alpha + \beta}{4} (x^2 - y^2)xy.$$

Il viendra

$$H_4 = y \frac{dF_4}{dx} + x \frac{dF_4}{dy} = \frac{3\beta - \alpha}{4} (x^2 + y^2)^2.$$

Si $3\beta = \alpha$, le premier membre de cette équation est nul, et l'on a satisfait à l'équation (4). Si $3\beta > \alpha$, le premier membre est positif, quels que soient x et y . Si $3\beta < \alpha$, le premier membre est négatif, quels que soient x et y .

Cela posé, on voit aisément que l'on peut faire deux hypothèses :

1° On peut supposer d'abord que l'on puisse déterminer $F_2, F_3, \dots, F_{q-1}$, de façon à satisfaire aux $q-2$ premières équations (1); mais qu'il soit impossible ensuite de déterminer F_q de façon à satisfaire à la

q — $i^{\text{ème}}$ équation (1). Dans ce cas, q est nécessairement pair, et nous déterminerons F_q , comme il vient d'être dit, de telle façon que

$$H_q - y \frac{dF_q}{dx} + x \frac{dF_q}{dy} = C_0 (x^2 + y^2)^{\frac{q}{2}}.$$

Nous poserons ensuite

$$F = F_2 + F_3 + \dots + F_{q-1} + F_q.$$

Je dis que, si k est une constante positive suffisamment petite, l'équation

$$F = k$$

représentera une courbe fermée qui sera un cycle sans contact.

En effet, si ρ est suffisamment petit (plus petit que ρ_0 , par exemple), la fonction F va en décroissant quand ρ décroît de ρ_0 à zéro, ω restant constant; car, si ρ est très petit, c'est $\frac{dF_2}{d\rho} = 2\rho$ qui donne son signe à $\frac{dF}{d\rho}$.

Soit k_0 la plus petite valeur que puisse prendre F le long du cercle $\rho = \rho_0$, et soit $k < k_0$.

Il est clair que $F = k$ sera une courbe fermée, ou plutôt que, parmi les branches dont se compose cette courbe algébrique, il y en a une qui est fermée, qui enveloppe l'origine et qui est intérieure au cercle $\rho = \rho_0$. C'est cette branche que nous envisagerons à l'exclusion de toutes les autres.

Maintenant, pour qu'il y eût contact entre ce cycle et une de nos trajectoires, il faudrait que l'on eût

$$\Phi = X \frac{dF}{dx} + Y \frac{dF}{dy} = 0.$$

Or le polynôme Φ , par hypothèse, n'a pas de terme de degré inférieur à q , et ses termes de degré q se réduisent à

$$- C_0 (x^2 + y^2)^{\frac{q}{2}}.$$

Si ρ est inférieur à une certaine limite (et nous pourrions toujours supposer que ρ_0 est inférieur à cette limite), ce sont ces termes de degré q qui donnent leur signe à Φ , et, comme ils sont toujours de même signe, Φ ne peut pas s'annuler.

Il ne peut donc y avoir de contact entre nos deux courbes

Ainsi l'intérieur de la courbe $F = k_0$ est sillonné par une infinité de cycles sans contact s'enveloppant mutuellement et enveloppant l'origine.

Supposons C_0 négatif pour fixer les idées, on aura, à l'intérieur du cercle $\rho = \rho_0$,

$$\frac{dF}{dt} = X \frac{dF}{dx} + Y \frac{dF}{dy} > 0.$$

Donc, lorsque t tendra vers $+\infty$, le point mobile ira en s'éloignant de l'origine jusqu'à ce qu'il soit sorti de la courbe $F = k_0$, et, une fois sorti de cette courbe, il n'y pourra plus rentrer. Lorsque t tendra vers $-\infty$, le point mobile se rapprochera indéfiniment et asymptotiquement de l'origine en décrivant une infinité de spires autour de ce point. La trajectoire est donc une spirale.

En d'autres termes, *il y aura instabilité, et l'origine sera un foyer.*

Si C_0 était positif, il suffirait de changer le signe de t , et l'on retomberait sur les mêmes résultats.

Soient, par exemple,

$$\frac{dx}{dt} = y - \frac{\beta x^3}{2}, \quad \frac{dy}{dt} = -x + \frac{2\alpha x^2 y - \beta y^3}{2}.$$

Nous ferons $F_2 = x^2 + y^2$, $F_3 = 0$, et la troisième équation (1) s'écrira

$$y \frac{dF_3}{dx} - x \frac{dF_3}{dy} = \beta x^4 - 2\alpha x^2 y^2 + \beta y^4 = H_4.$$

Nous avons vu qu'il est impossible, en général, de satisfaire à cette équation. Nous prendrons

$$F_4 = \frac{\alpha + \beta}{4} (y^2 - x^2)xy,$$

et il viendra, comme nous l'avons vu,

$$H_4 = y \frac{dF_4}{dx} + x \frac{dF_4}{dy} = \frac{3\beta - \alpha}{2} (x^2 + y^2)^2.$$

D'après ce que nous venons de voir, si $3\beta - \alpha$ n'est pas nul, les courbes

$$F_2 + F_4 = k$$

seront des cycles sans contact, pourvu que k soit suffisamment petit. Il y aura instabilité, et l'origine sera un foyer.

Supposons maintenant que $3\beta = \alpha$. La troisième équation (1) sera alors satisfaite. Nous prendrons $F_5 = 0$, et la cinquième équation (1) s'écrit

$$y \frac{dF_6}{dx} + x \frac{dF_6}{dy} = H_6.$$

On trouve, d'ailleurs, en faisant $\beta = 1$, $\alpha = 3$,

$$H_6 = \frac{1}{2}x^3y - 9x^3y^3 + \frac{1}{2}y^3x,$$

et l'on voit qu'il est possible de satisfaire à la cinquième équation (1) en faisant

$$F_6 = \frac{1}{4}x^6 + 3x^4y^2 + \frac{1}{4}x^2y^4.$$

Nous prendrons ensuite $F_7 = 0$, et la septième équation (1) s'écrit

$$y \frac{dF_8}{dx} + x \frac{dF_8}{dy} = H_8,$$

où

$$8H_8 = 30x^8 - 96x^6y^2 - 42x^4y^4 + 12x^2y^6$$

ou

$$\frac{1}{3}H_8 = 5x^8 - 16x^6y^2 - 7x^4y^4 + 2x^2y^6.$$

Posons $x = \rho \cos \omega$, $y = \rho \sin \omega$; développons ensuite H_8 suivant les sinus et cosinus des multiples de ω , et voyons si le terme tout connu

$$\frac{1}{3}C_0\rho^8$$

est nul. On trouve

$$\frac{1}{3} \frac{H_3}{\beta^3} = 12 \cos^8 \omega + 4 \cos^6 \omega - 13 \cos^4 \omega + 2 \cos^2 \omega;$$

d'où

$$\frac{1}{3} C_0 = \frac{12 \cdot 33}{128} + \frac{4 \cdot 5}{16} - \frac{13 \cdot 6}{16} + \frac{2}{2} = \frac{81}{128}.$$

Ainsi il est impossible de satisfaire à la septième équation (1); l'origine est donc encore un foyer.

On doit conclure de là que, si le mouvement d'un point mobile est défini par les équations

$$\frac{dx}{dt} = y - \frac{\beta x^3}{2}, \quad \frac{dy}{dt} = -x + \frac{2\alpha x^2 y - \beta y^3}{2},$$

la trajectoire de ce point sera toujours instable, quels que soient α et β .

Ainsi l'on cherchera à résoudre successivement toutes les équations (1), et, dès qu'on se trouvera arrêté, on sera certain que la trajectoire est instable.

2° Mais il peut arriver aussi qu'on ne soit jamais arrêté, ce qui exige une infinité de conditions. Ces conditions sont évidemment nécessaires, pour que la trajectoire soit stable, ou pour que l'origine soit un centre. Sont-elles suffisantes? C'est ce que nous allons examiner.

Formons successivement, à l'aide des équations (1), les polynômes $F_2, F_3, \dots, F_q$, et considérons la série infinie

$$F = F_2 + F_3 + \dots + F_q + \dots$$

Si elle est convergente, il n'y a pas de difficulté, car elle satisfera à l'équation

$$\frac{dF}{dx} X + \frac{dF}{dy} Y = 0.$$

Les courbes $F = k$ seront donc les trajectoires du point mobile, et ce seront des courbes fermées si k est suffisamment petit.

Il reste à examiner si la série F converge.

Posons

$$X = R \cos \omega - \rho \Omega \sin \omega, \quad Y = R \sin \omega + \rho \Omega \cos \omega.$$

L'équation précédente deviendra

$$\frac{dF}{d\rho} R + \frac{dF}{d\omega} \Omega = 0.$$

Nous pourrions développer la fonction $-\frac{R}{\Omega}$ suivant les puissances croissantes de ρ , le développement commençant par un terme en ρ^2 : nous écrirons

$$-\frac{R}{\Omega} = \rho^2 \zeta_2 + \rho^3 \zeta_3 + \dots,$$

$\zeta_2, \zeta_3, \dots$ étant des fonctions de ω . L'équation précédente deviendra

$$\frac{dF}{d\omega} = \frac{dF}{d\rho} (\rho^2 \zeta_2 + \rho^3 \zeta_3 + \dots),$$

et, si l'on pose

$$F_q = \rho^q z_q, \quad z'_q = \frac{dz_q}{d\omega},$$

les z_q étant des fonctions de ω .

Alors on déterminera successivement les z_q à l'aide des équations suivantes qui remplaceront les équations (1) :

$$(1 \text{ bis}) \left\{ \begin{array}{l} z_2 = 1, \\ z'_3 = 2z_2 \zeta_2, \\ z'_4 = 3z_3 \zeta_2 + 2z_2 \zeta_3, \\ z'_5 = 4z_4 \zeta_2 + 3z_3 \zeta_3 + 2z_2 \zeta_4, \\ \dots\dots\dots \\ z'_q = (q-1)z_{q-1} \zeta_2 + (q-2)z_{q-2} \zeta_3 + \dots + 3z_3 \zeta_{q-2} + 2z_2 \zeta_{q-1}, \\ \dots\dots\dots \end{array} \right.$$

S'il est possible de satisfaire aux équations (1) avec des polynômes entiers en x et en y , il sera possible de satisfaire aux équations (1 bis) avec des fonctions purement trigonométriques de ω (c'est-à-dire avec des polynômes en $\cos \omega$ et $\sin \omega$).

Si, au contraire, il n'est pas possible de satisfaire aux équations (1) (c'est-à-dire si les quantités C_0 ne sont pas toutes nulles), on pourra

néanmoins résoudre les équations (1 bis) et calculer successivement les fonctions z'_q ; mais ces fonctions, au lieu de ne contenir que des termes trigonométriques, contiendront des termes où ω entrera, en dehors des signes sinus et cosinus, soit à la première puissance, soit à une puissance supérieure.

Il est impossible de n'être pas frappé de l'analogie des termes ainsi introduits avec les termes que les astronomes appellent *séculaires*. Il y a, cependant, une différence essentielle qu'il importe de remarquer. Quand, dans les méthodes habituelles de la Mécanique céleste, on rencontre un terme séculaire, il n'est pas permis, pour cela, de conclure à l'instabilité de l'orbite; car il peut se faire, ou bien que la série soit divergente, ou bien que le terme ainsi obtenu ne soit que le premier terme d'un développement dont la somme reste toujours finie. C'est ainsi que $z\omega$ peut être le premier terme du développement

$$\sin z\omega = z\omega - \frac{z^3\omega^3}{6} + \frac{z^5\omega^5}{120} - \dots$$

Il n'en est pas de même dans le cas qui nous occupe et avec la méthode que je viens d'exposer. Si, dans la suite des calculs, on rencontre un terme séculaire, on pourra conclure immédiatement à l'instabilité; il n'est pas même nécessaire, pour cela, que la série

$$F = \rho^2 z_2 + \rho^3 z_3 + \dots$$

soit convergente.

Nous pouvons poser la question de la convergence même dans le cas où les fonctions z_q contiennent des termes séculaires; car, bien que cette question présente alors beaucoup moins d'intérêt, il est avantageux, pour arriver plus facilement à la solution, de se débarrasser des restrictions inutiles.

Commençons par dire quelques mots des trois cas simples suivants :

$$-\frac{R}{\Omega} = \rho^2 \frac{dz_1}{d\omega}, \quad -\frac{R}{\Omega} = (\rho^2 + \rho^3) \frac{dz_1}{d\omega}, \quad -\frac{R}{\Omega} = (\rho^2 + \rho^4) \frac{dz_1}{d\omega}.$$

On trouve alors, pour les intégrales générales des équations du mouve-

ment, en appelant k une constante d'intégration,

$$\frac{1}{z} - \varphi = k, \quad \frac{1}{z} + 1, \frac{z}{z+1} - \varphi = k, \quad \frac{1}{z} + \arctan \rho - \varphi = k.$$

Cherchons à former F .

Dans le premier cas, on trouve aisément

$$F = \frac{z^2}{(1 - z\varphi)^2}.$$

Dans le troisième cas, si nous posons

$$\frac{1}{1 + \frac{z}{z+1} \arctan \varphi} = \psi,$$

le premier membre de cette égalité sera une fonction holomorphe de z (pour $z = 0$) qui s'annule avec z , mais dont la dérivée ne s'annule pas avec z . On en déduira, d'après un théorème connu,

$$z = \psi' \xi,$$

ψ étant une fonction holomorphe de ξ . On aura alors

$$F = \left| \psi \left(\frac{1}{1 + \frac{z}{z+1} \arctan \varphi} - \psi \right) \right|^2.$$

Cette fonction, comme dans le premier cas, est holomorphe en z et z , pourvu que ces variables soient suffisamment petites.

Dans le deuxième cas, on ne peut appliquer ce procédé, parce que la fonction

$$\frac{1}{\frac{1}{z} + 1, \frac{z}{z+1}}$$

n'est pas holomorphe. Posons alors

$$F = H_0 + H_1 z + \frac{H_2 z^2}{1,2} + \dots + \frac{H_n z^n}{1,2,\dots,n} + \dots,$$

en développant F non plus suivant les puissances de ρ , mais suivant

celles de φ . On déterminera ensuite les fonctions H successivement à l'aide des équations suivantes :

$$(1^{er}) \quad \left\{ \begin{array}{l} H_0 = \rho^2, \\ H_1 = (\rho^2 + \rho^2) \frac{dH_0}{d\varphi}, \\ H_2 = (\rho^2 + \rho^2) \frac{dH_1}{d\varphi}, \\ \dots\dots\dots \\ H_n = (\rho^2 + \rho^2) \frac{dH_{n-1}}{d\varphi}, \\ \dots\dots\dots \end{array} \right.$$

Les fonctions H ainsi définies sont des polynômes entiers en ρ , et il est aisé de voir que tous les coefficients sont positifs, que le degré de H_n est $2(n+1)$, et que ce polynôme H_n ne contient pas de terme de degré plus petit que $n+2$.

Soient S_n la somme des coefficients de H_n et S'_n celle des coefficients de sa dérivée $\frac{dH_n}{d\varphi}$; H_n étant de degré $2(n+1)$; on aura

$$S'_n < 2S_n(n+1).$$

D'ailleurs, les formules (1^{er}) nous donnent

$$S_{n+1} = 2S'_n;$$

d'où

$$S_{n+1} < 4S_n(n+1)$$

et

$$S_{n+1} < 4^{n+1}(n+1)!$$

Supposons ρ positif et plus petit que 1, et envisageons le terme général de la série qui définit F ; on aura

$$\left| \frac{H_n \varphi^n}{n!} \right| < (4\rho\varphi)^n.$$

Si donc $\rho\varphi$ est positif et plus petit que $\frac{1}{4}$, la série est convergente et, comme tous ses termes sont positifs, absolument convergente.

La conclusion, c'est que F est une fonction holomorphe de ρ et de φ , pourvu que

$$|\rho| < 1, \quad |\rho\varphi| < \frac{1}{2}.$$

Il était aisé de prévoir ce résultat. Posons, en effet,

$$(5) \quad \frac{1}{\rho} + L \frac{\rho}{\rho+1} - \varphi = \frac{1}{\zeta} + L \frac{\zeta}{\zeta+1}.$$

Je dis que ζ est une fonction holomorphe de ρ et de φ dans le voisinage du point $\rho = \varphi = 0$. Pour cela, il faut démontrer deux choses :

1° Que ζ tend vers zéro, toutes les fois que ρ et φ tendent simultanément vers zéro.

En effet, si ρ et φ tendent vers zéro, les deux membres de l'équation (5) croîtront indéfiniment. Or, pour que

$$\frac{1}{\zeta} + L \frac{\zeta}{\zeta+1}$$

croisse indéfiniment, il faut que ζ tende vers zéro ou vers -1 . Mais, si la valeur initiale de ζ est suffisamment voisine de zéro, il faudra que ζ tende vers zéro et non pas vers -1 . Il suffira de le vérifier, ce qui est facile, lorsque, φ et l'argument de ρ restant constants, le module de ρ tend vers zéro. Cela sera suffisant, parce que nous allons voir un peu plus loin que ζ est une fonction *uniforme* de ρ et de φ .

2° Il faut démontrer ensuite que ζ revient à la même valeur quand ρ décrit dans son plan, et ω dans le sien, un contour suffisamment petit enveloppant le point zéro. Or, dans ces conditions, le premier et, par conséquent, le second membre de l'équation (5) augmenteront d'un multiple de $2i\pi$, ce qui fera décrire au point ζ un contour fermé enveloppant le point zéro.

Donc ζ , et, par conséquent,

$$F = \zeta^2$$

sont une fonction holomorphe de ρ et de φ si ces variables sont assez petites.

Supposons maintenant

$$-\frac{R}{\Omega} = P(\rho) = \rho^2 + \beta\rho^3 + \gamma\rho^4 + \dots,$$

$P(\rho)$ étant une série ordonnée suivant les puissances de ρ et convergente, pourvu que ρ soit suffisamment petit.

L'intégrale générale des équations différentielles sera

$$\omega + \int \frac{d\rho}{P(\rho)} = \text{const.}$$

On trouvera, d'ailleurs,

$$\int \frac{d\rho}{P(\rho)} = -\frac{1}{\rho} - \beta L\rho - Q(\rho),$$

$Q(\rho)$ étant une fonction de ρ holomorphe pour $\rho = 0$.

Posons maintenant

$$\text{5 bis)} \quad \frac{1}{\rho} + \beta L\rho + Q(\rho) - \omega = \frac{1}{\xi} + \beta L\xi + Q(\xi).$$

Nous considérerons, parmi les fonctions ξ qui satisfont à cette équation, celle qui se réduit à ρ pour $\omega = 0$. Je dis que ce sera une fonction holomorphe de ρ et de ω , si ces variables sont suffisamment petites.

Pour cela, il faut faire voir que, si ρ et ω sont assez petits, ξ est une fonction uniforme de ρ et de ω qui tend vers zéro quand ces variables tendent simultanément vers zéro. Le raisonnement serait absolument le même que dans le cas précédent. Il en résulte que

$$F = \xi^2$$

est une fonction holomorphe de ρ et de ω .

Il est, d'ailleurs, aisé de trouver les coefficients du développement de F suivant les puissances de ρ et de ω . Écrivons, en effet,

$$F = H_0 + H_1\omega + \frac{H_2\omega^2}{2} + \dots + \frac{H_n\omega^n}{n!} \dots,$$

$$P(\rho) = \Sigma a_p \rho^p,$$

$$H_n = \Sigma h_{np} \rho^p.$$

Nous supposerons, ce qui est utile pour notre objet, que tous les a_p sont positifs, et nous pourrions trouver deux nombres μ et α , tels que

$$a_p < \mu \alpha^p.$$

Les H nous seront donnés par les équations

$$\begin{aligned} H_0 &= \rho^2, \\ H_1 &= P(\rho) \frac{dH_0}{d\rho}, \\ &\dots\dots\dots \\ H_{n+1} &= P(\rho) \frac{dH_n}{d\rho}. \end{aligned}$$

Nous adopterons la notation suivante : l'inégalité

$$f(\rho) \leq \varphi(\rho),$$

avec un double signe d'inégalité, signifiera (lorsque les coefficients des fonctions f et φ développées suivant les puissances de ρ sont positifs, ce que nous supposerons) que chaque coefficient de f est plus petit que le coefficient correspondant de φ . Nous pourrions écrire alors

$$P(\rho) \leq \frac{\mu}{1 - a\rho}.$$

Je dis qu'on pourra toujours trouver un nombre M_n , tel que

$$H_n(\rho) \leq \frac{M_n n!}{(1 - a\rho)^{2n+1}}.$$

Supposons, en effet, que cela soit vrai pour H_n , je dis que cela sera vrai pour H_{n+1} . Il viendra, dans cette hypothèse,

$$\frac{dH_n}{d\rho} \leq \frac{M_n a n! (2n+1)}{(1 - a\rho)^{2n+2}} \leq \frac{2a M_n (n+1)!}{(1 - a\rho)^{2n+2}};$$

d'où

$$H_{n+1} = P \frac{dH_n}{d\rho} \leq \frac{2a \mu M_n (n+1)!}{(1 - a\rho)^{2n+3}};$$

d'où

$$M_{n+1} = 2a\mu M_n, \quad M_n = M_0 (2a\mu)^n.$$

Il vient alors

$$F = \sum \frac{H_n \omega^n}{n!} \leq \sum \frac{M_0 (2a\mu \omega)^n}{(1 - a\rho)^{2n+1}} = \frac{M_0}{1 - a\rho} \frac{1}{1 - \frac{2a\mu \omega}{(1 - a\rho)^2}}.$$

On conclut de cette inégalité que la série F converge, pourvu que, par exemple,

$$|\rho| < \frac{1}{2a}, \quad |\omega| < \frac{1}{8a\mu}.$$

Mais il est aisé de voir que le développement de H_0 commence par un terme en ρ^2 , celui de H_1 par un terme en ρ^3 , ..., celui de H_n par un terme en ρ^{n+2} . Donc la fonction F est une fonction holomorphe, non seulement de ρ et de ω , mais de ρ et de $\rho\omega$. La série F convergera donc, pourvu que

$$|\rho| < \frac{1}{2a}, \quad |\rho\omega| < \frac{1}{16a^2\mu}.$$

Donc cette série sera convergente pour toutes les valeurs de ω comprises entre zéro et 2π , pourvu que

$$|\rho| < \frac{1}{2a}, \quad |\rho| < \frac{1}{32a^2\mu\pi}.$$

Voici comment ce qui précède se rattache aux principes exposés par MM. Briot et Bouquet dans le XXXVI^e Cahier du *Journal de l'Ecole Polytechnique*.

Considérons ω comme une constante; nous aurons, entre ζ et ρ , l'équation différentielle

$$\frac{d\zeta}{P(\zeta)} = \frac{d\rho}{P(\rho)}$$

ou

$$\frac{d\zeta}{d\rho} = \frac{\zeta^2}{\rho^2} \frac{1 + \beta\zeta + \dots}{1 + \beta\rho + \dots} = \frac{\zeta^2}{\rho^2} [1 + \beta(\zeta - \rho) + \psi],$$

ψ représentant un ensemble de termes de degré au moins égal à 2 en ζ et en ρ . Posons

$$\zeta = \rho(1 + v).$$

Il viendra

$$\rho \frac{dv}{d\rho} = v + v^2 + \beta(1 + v)^2 \rho v + (1 + v)^2 \psi,$$

la fonction ψ contenant ρ^2 en facteur.

C'est là un type d'équations qui, d'après un théorème de MM. Briot

et Bouquet, admet une infinité d'intégrales holomorphes, s'annulant avec ρ .

Donc ζ est fonction holomorphe de ρ .

C. Q. F. D.

Passons maintenant au cas général.

Il importe d'abord de rappeler et de préciser le sens de la notation $<$ déjà employée plus haut. Quand j'écrirai dans ce qui va suivre :

$$f(\rho, \omega) < \varphi(\rho, \omega),$$

je regarderai les deux fonctions f et φ comme développées suivant les puissances croissantes de ρ . Les coefficients des deux développements seront des fonctions de ω que je regarderai momentanément comme une constante et que je supposerai réelle et comprise entre zéro et 2π .

L'inégalité signifiera alors que, pour toutes les valeurs réelles de ω comprises entre zéro et 2π , tous les coefficients du développement de φ sont positifs et plus grands en valeur absolue que les coefficients correspondants du développement de f .

Soient

$$\begin{aligned} R &= \rho^2 + R_3 \rho^3 + \dots + R_p \rho^p, \\ \Omega &= 1 + \Omega_1 \rho + \Omega_2 \rho^2 + \dots + \Omega_q \rho^q. \end{aligned}$$

Les coefficients R_i et Ω_i des deux polynômes R et Ω seront des fonctions trigonométriques de ω . Supposons que toutes ces fonctions trigonométriques restent constamment inférieures en valeur absolue à une certaine quantité positive L . Il viendra

$$(6) \quad -\frac{R}{\Omega} < \frac{(\rho^2 + \rho^3 + \dots + \rho^p)L}{1 - (\rho + \rho^2 + \dots + \rho^q)L} < \frac{\rho^2 L}{1 - \rho(L+1)}.$$

La fonction $\frac{\rho^2 L}{1 - \rho(L+1)}$ qui ne dépend que de ρ est développable suivant les puissances de ρ , pourvu que ρ soit suffisamment petit.

Il en résulte qu'il existe une série F , ordonnée suivant les puissances croissantes de ρ et de ω , convergente pour toutes les valeurs réelles de ω comprises entre zéro et 2π , pourvu que ρ soit suffisamment petit, et satisfaisant à l'égalité

$$\frac{dF_1}{d\omega} = \frac{dF_1}{d\rho} \frac{\rho^2 L}{1 - \rho(L+1)}.$$

De plus F_1 se réduit à ρ^2 pour $\omega = 0$. Posons, comme plus haut,

$$\begin{aligned} F &= z_2 \rho^2 + z_3 \rho^3 + \dots + z_q \rho^q + \dots, \\ F_1 &= u_2 \rho^2 + u_3 \rho^3 + \dots + u_q \rho^q + \dots \end{aligned}$$

Les fonctions z_q seront définies par les équations (1 bis) et les fonctions u_q par les équations analogues

$$\begin{aligned} u_2 &= 1, \\ (7) \quad u'_q &= (q-1)u_{q-1}\theta'_2 + (q-2)u_{q-2}\theta'_3 + \dots + 2u_2\theta'_{q-1}, \end{aligned}$$

où

$$\theta'_{q+2} = (L+1)^q L.$$

Cela posé, je dis qu'on aura constamment, ω étant réel et plus petit que 2π ,

$$(8) \quad |z_q| < u_q.$$

Pour cela, je vais supposer que l'inégalité (8) a lieu pour $z_2, z_3, \dots, z_{q-1}$, et démontrer qu'elle a encore lieu pour z_q .

En effet, s'il en est ainsi, on a, en comparant les relations (1 bis), (6) et (7),

$$(9) \quad |z'_q| < u'_q.$$

La fonction z_q n'est pas entièrement déterminée par les équations (1 bis); ces équations ne nous donnent en effet que la dérivée de z_q en fonction de $z_2, z_3, \dots, z_{q-1}$. Il en résulte que z_q n'est connue qu'à une constante d'intégration près. Nous disposerons de cette constante de telle façon que z_q s'annule avec ω .

Dans ces conditions, l'inégalité (9) entraîne l'inégalité

$$(8) \quad |z_q| < u_q, \quad \text{C. Q. F. D.}$$

Donc on a

$$F \leq F_1.$$

et, comme pour les petites valeurs de ρ , F , est convergent quand ω varie de zéro à 2π , il en sera de même de F .

Mais, d'après la façon dont nous avons déterminé les z_q , ce sont des fonctions trigonométriques de ω (dans le cas où tous les C_0 sont nuls). Si donc F converge pour les valeurs de ω comprises entre zéro et 2π , cette série devra converger pour toutes les valeurs réelles de ω .

Il est à remarquer que, si, partant de la série F telle que nous venons de la définir, on repasse des coordonnées polaires ρ et ω aux coordonnées rectilignes x et y , F ne sera plus une série ordonnée suivant les puissances croissantes de x et de y . Cela tient à la façon dont nous avons disposé des constantes d'intégration, de façon que z_q s'annule avec ω .

Il pourra se faire alors, par exemple, que l'on ait

$$F = \rho^2 + \rho^3(1 - \cos \omega) + \dots$$

ce qui donne

$$F = x^2 + y^2 + (x^2 + y^2)^{\frac{3}{2}} - x(x^2 + y^2) + \dots$$

En effet, si l'on tient à ce que F soit une fonction holomorphe de x et de y , on peut disposer de la constante d'intégration (que nous avons appelée plus haut A_0) lorsque q est pair, mais on n'en peut plus disposer quand q est impair. Il ne sera donc pas possible en général de s'arranger pour que z_q s'annule avec ω .

Cela n'a d'ailleurs que peu d'importance au point de vue qui nous occupe, mais il est aisé de tourner la difficulté. On a

$$\frac{dF}{d\omega} = -\frac{R}{\Omega} \frac{dF}{d\rho}.$$

Soit F_2 ce que devient F quand on y change ρ en $-\rho$, et ω en $\omega + \pi$. On aura encore, comme il est aisé de le vérifier,

$$\frac{dF_2}{d\omega} = -\frac{R}{\Omega} \frac{dF_2}{d\rho},$$

car $-\frac{R}{\Omega}$ change de signe quand on y change ρ et $-\rho$, et ω en $\omega + \pi$.

On aura donc encore

$$\frac{d(F + F_2)}{d\omega} = -\frac{R}{\Omega} \frac{d(F + F_2)}{d\rho},$$

et la série $F + F_2$ sera convergente pour toutes les valeurs réelles de ω , pourvu que ρ soit suffisamment petit. Si d'ailleurs on repasse aux coordonnées rectilignes x et y , $F + F_2$ sera une fonction holomorphe de x et de y . On voit donc qu'il existe toujours, si tous les C_0 sont nuls, une série F convergente, ordonnée suivant les puissances de x et de y et satisfaisant à l'équation

$$X \frac{dF}{dx} + Y \frac{dF}{dy} = 0.$$

En résumé, *pour que l'origine soit un centre, c'est-à-dire pour que la trajectoire du point mobile soit stable, il faut et il suffit que toutes les quantités que nous avons appelées C_0 soient nulles à la fois.*

Toutefois, il sera souvent difficile de reconnaître si ces conditions, en nombre infini, sont remplies à la fois. Il y a donc intérêt à signaler des cas où l'on est certain d'avance que tous les C_0 sont nuls.

Je ne signalerai que le plus simple d'entre eux.

Supposons que, quand on change y en $-y$, sans changer x , X se change en $-X$ et que Y ne change pas. Je dis que les trajectoires du point mobile seront des courbes fermées symétriques par rapport à l'axe des x .

En effet, partons, pour $t = 0$, d'un point initial situé sur l'axe des x . La vitesse initiale du point mobile sera, d'après les hypothèses faites, perpendiculaire à cet axe. Si l'on change t en $-t$, y en $-y$, x en $-x$, les équations du mouvement et ses conditions initiales ne changent pas. Donc le point mobile occupera aux temps t en $-t$ deux points du plan symétrique, par rapport à l'axe des x . D'après la forme des équations, si le point de départ de notre mobile est suffisamment voisin de l'origine, il arrivera à une époque t_0 où sa trajectoire viendra recouper l'axe des x . Ainsi, aux deux époques t_0 et $-t_0$, le point mobile occupera un même point de l'axe des x . Sa trajectoire sera donc une courbe fermée, et il résulte de ce qui précède qu'elle doit

être symétrique par rapport à l'axe des x . Donc on est certain d'avance que tous les C_0 sont nuls.

On rencontre un exemple du cas que nous venons de signaler dans un problème astronomique sur lequel M. Tisserand a bien voulu appeler mon attention. Delaunay a rencontré dans sa théorie de la Lune les équations suivantes :

$$\begin{aligned}\frac{de}{dt} &= M(1 + M_1 e^2 + M_2 e^4) \sin \zeta, \\ \frac{d\theta}{dt} &= N(1 + N_1 e^2 + N_2 e^4 + N_3 e^6) + \frac{M}{e}(1 + P_1 e^2 + P_2 e^4) \cos \zeta\end{aligned}$$

On suppose que, pour $t = 0$, e est très petit et que, le coefficient M étant de l'ordre du carré de cette petite quantité, les autres coefficients sont finis (*Mémoires de l'Académie des Sciences*, t. XXVIII, p. 107).

Delaunay donne les expressions suivantes

$$\begin{aligned}e \cos \zeta &= \sum A_i \cos i(zt + c), \\ e \sin \zeta &= \sum B_i \sin i(zt + c);\end{aligned}$$

mais il ne traite pas la question de la convergence et de la possibilité du développement.

Les équations sont de même forme que celles que nous étudions. Posons, en effet,

$$e \cos \zeta = x, \quad e \sin \zeta = y,$$

il viendra

$$\begin{aligned}\frac{dx}{dt} &= Mx\mathcal{Y}(M_1 - P_1 + M_2 e^2 - P_2 e^2) \\ &\quad - N\mathcal{Y}(1 + N_1 e^2 + N_2 e^4 + N_3 e^6) = X, \\ \frac{dy}{dt} &= M + M\mathcal{Y}^2(M_1 + M_2 e^2) \\ &\quad + Mx^2(P_1 + P_2 e^2) + Nx(1 + N_1 e^2 + N_2 e^4 + N_3 e^6) = Y,\end{aligned}$$

ou

$$e^2 = x^2 + y^2;$$

X et Y s'annulent, pour $y = 0$,

$$(10) \quad M(1 + P_1 x^2 + P_2 x^4) + Nx(1 + N_1 x^2 + N_2 x^4 + N_3 x^6) = 0.$$

En vertu des hypothèses faites sur les coefficients, l'équation (10) est satisfaite pour $x = x_1$, x_1 étant une quantité très petite de l'ordre de M .

Le point $x = x_1$, $y = 0$ est un centre; car X change de signe et Y ne change pas quand on change y en $-y$. On est donc certain d'avance que toutes les quantités que nous avons appelées C_0 sont nulles à la fois.

Il en résulte que x et y sont des fonctions périodiques du temps t , qui peuvent être représentées par des séries de la forme obtenue par Delaunay.

Si l'on a reconnu d'une manière ou d'une autre que toutes les quantités C_0 sont nulles, on est certain qu'il y a autour du centre une certaine région du plan R qui est sillonnée par des courbes fermées ou cycles enveloppant le centre et qui sont les trajectoires du point mobile dans la région considérée. Au delà de la région R , les trajectoires seront en général des spirales. Cette région sera limitée par un certain cycle frontière qui sera la dernière trajectoire fermée. Je dis que ce cycle frontière doit passer par un point singulier.

En effet nous pourrions toujours tracer un arc sans contact venant couper ce cycle frontière, ainsi que les trajectoires fermées qui en sont très voisines et se prolongeant au delà de ce cycle frontière et en dehors de la région R . Nous définirons la position d'un point sur cet arc à l'aide d'un paramètre t qui sera, par exemple, nul sur le cycle frontière, négatif à l'intérieur de la région R et positif à l'extérieur de cette région.

Reportons-nous maintenant au Chapitre V (II^e Partie) et à ce que nous avons appelé la loi de conséquence

$$t_1 = \varphi_1(t_0).$$

Pour les valeurs négatives de t_0 , on est à l'intérieur de R ; les trajectoires sont fermées et l'on a

$$\varphi_1(t_0) = t_0.$$

Au contraire, pour les valeurs positives de t_0 , on est hors de R ; les

trajectoires ne sont plus fermées, et l'on a

$$\varphi_1(t_0) \geq t_0.$$

Il est donc impossible que la fonction φ_1 soit holomorphe pour $t_0 = 0$. Donc, en vertu du théorème XIII (p. 255), le cycle frontière doit aller passer par un point singulier.

CHAPITRE XII.

ÉQUATIONS DE DEGRÉ SUPÉRIEUR.

Nous allons étudier maintenant les équations différentielles du premier ordre et de degré supérieur, c'est-à-dire les équations de la forme

$$(1) \quad F\left(x, y, \frac{dy}{dx}\right) = 0,$$

F étant un polynôme entier en x, y et $\frac{dy}{dx}$.

Voici le mode de représentation géométrique que nous adopterons. Nous pourrions d'abord envisager la surface

$$(2) \quad F(x, y, z) = 0,$$

et exprimer les équations du mouvement du point mobile sur cette surface, de la façon suivante

$$(2 \text{ bis}) \quad \frac{dx}{dt} = -\frac{dF}{dz}, \quad \frac{dy}{dt} = -z \frac{dF}{dz}, \quad \frac{dz}{dt} = \frac{dF}{dx} + z \frac{dF}{dy},$$

de telle sorte que $\frac{dx}{dt}, \frac{dy}{dt}, \frac{dz}{dt}$ sont égaux à des polynômes entiers en x, y, z .

C'est là le mode de représentation le plus simple, toutes les fois que la surface $F(x, y, z)$ n'a pas de nappes infinies ou de singularités gênantes. Mais il peut être avantageux, dans certains cas, d'employer un mode de représentation plus général.

Posons

$$(3) \quad \xi = \varphi_1(x, y, z), \quad \eta = \varphi_2(x, y, z), \quad \zeta = \varphi_3(x, y, z),$$

les fonctions φ_1 , φ_2 et φ_3 étant rationnelles en x , y , z .

Sauf des cas exceptionnels que nous laisserons de côté, on déduira des équations (2) et (3) les équations

$$(4) \quad F_1(\xi, \eta, \zeta) = 0$$

et

$$(5) \quad x = \theta_1(\xi, \eta, \zeta), \quad y = \theta_2(\xi, \eta, \zeta), \quad z = \theta_3(\xi, \eta, \zeta),$$

F_1 étant un polynôme entier et θ_1 , θ_2 , θ_3 des fonctions rationnelles. Les deux surfaces (2) et (4) se correspondront alors point par point par une transformation birationnelle, et l'on aura

$$\frac{d\xi}{\psi_1(\xi, \eta, \zeta)} = \frac{d\eta}{\psi_2(\xi, \eta, \zeta)} = \frac{d\zeta}{\psi_3(\xi, \eta, \zeta)},$$

les fonctions ψ_1 , ψ_2 et ψ_3 étant rationnelles.

On pourra disposer de ce qu'il y a d'indéterminé dans la transformation birationnelle (3) pour que la surface (4) n'ait pas de nappes infinies et aussi pour atteindre différents autres buts, par exemple pour faire disparaître des singularités gênantes. Voici donc comment nous nous poserons le problème des équations différentielles de degré supérieur.

On donne une surface S ayant pour équation

$$F(x, y, z) = 0$$

et n'ayant pas de nappes infinies; et l'on demande d'étudier le mouvement d'un point mobile sur cette surface, les équations du mouvement étant

$$\frac{dx}{dt} = X, \quad \frac{dy}{dt} = Y, \quad \frac{dz}{dt} = Z,$$

où X , Y et Z sont des polynômes entiers en x , y , z avec la condition

$$\frac{dF}{dx}X + \frac{dF}{dy}Y + \frac{dF}{dz}Z = MF.$$

Étudions d'abord les trajectoires du point mobile dans le voisinage d'un point M de la surface. Nous distinguerons le cas où le point M est un point ordinaire de la surface S (tout en pouvant être un point singulier pour les équations différentielles) et le cas où le point M est un point singulier de la surface S .

Dans le premier cas, on pourra exprimer, dans le voisinage du point M , x , y et z par des fonctions holomorphes de deux paramètres u et v et de telle façon que les trois déterminants fonctionnels

$$\frac{\partial(x, y)}{\partial(u, v)}, \quad \frac{\partial(y, z)}{\partial(u, v)} \quad \text{et} \quad \frac{\partial(x, z)}{\partial(u, v)}$$

ne soient pas nuls à la fois.

On pourra écrire alors

$$\frac{du}{dt} = U, \quad \frac{dv}{dt} = V,$$

U et V étant des fonctions holomorphes en u et en v . On est alors ramené à l'étude des courbes planes définies par une équation différentielle du premier ordre et du premier degré; car, dans le voisinage du point considéré, les fonctions U et V ont tous les caractères des polynômes entiers.

Si le point considéré est un point ordinaire, il passe par ce point une trajectoire et une seule.

Si c'est un point singulier de l'équation différentielle, c'est-à-dire si $U = V = 0$, ce peut être un *col*, un *foyer*, un *nœud* ou un *centre*, présentant les mêmes propriétés que les points de même nom définis dans la I^{re} Partie.

Dans le cas des équations (2) et (2 bis) les points singuliers sont donnés par les équations

$$F = 0, \quad \frac{dF}{dz} = \frac{dF}{dx} + z \frac{dF}{dy} = 0.$$

Supposons maintenant que le point envisagé soit un point singulier pour la surface S elle-même, c'est-à-dire que l'on ait

$$\frac{dF}{dx} = \frac{dF}{dy} = \frac{dF}{dz} = 0.$$

Ce cas se ramène au précédent. Supposons d'abord, par exemple, que la surface S présente une courbe double, que le point considéré soit un point de cette courbe double, et que les deux plans tangents en ce point soient distincts. Alors on pourra exprimer, dans le voisinage du point envisagé, x , y et z en fonctions holomorphes de deux paramètres u et v , et cela de deux manières, l'une des manières se rapportant à l'une des nappes de la surface qui passent par la courbe double, et la seconde manière à l'autre nappe. On retombe donc sur le cas précédent.

De même, supposons que le point considéré, que nous pourrions prendre pour origine des coordonnées, soit un point conique du second ordre. Soit, par exemple,

$$F = F_2 + F_3 + \dots + F_n,$$

F_i étant un polynôme homogène de degré i en x , y , z . Considérons la portion de cette surface qui est voisine de l'origine, c'est-à-dire du point conique. Je dis que nous pourrions, par une transformation birationnelle, transformer cette portion de surface en une autre qui n'aura pas de point singulier. Il suffit pour cela de poser

$$\xi = \frac{x}{z}(1-z), \quad \eta = \frac{y}{z}(1-z), \quad \zeta = 1-z,$$

d'où

$$x = \frac{\xi}{\zeta}(1-\zeta), \quad y = \frac{\eta}{\zeta}(1-\zeta), \quad z = 1-\zeta.$$

Cette transformation birationnelle est donc réciproque et elle a pour effet de transformer la surface F dans la suivante :

$$z^{n-2}F_2 + z^{n-3}(1-z)F_3 + z^{n-4}(1-z)^2F_4 + \dots + (1-z)^nF_n = 0.$$

Cette surface transformée est coupée par le plan $z = 1$ suivant la conique

$$F_2(x, y, 1) = 0,$$

et elle ne présente pas de point singulier le long de cette conique. D'ailleurs, la portion de la surface S , voisine du point conique, devient, après la transformation, la portion de la surface transformée voisine de cette conique, c'est-à-dire une portion de surface dépourvue de point singulier.

On est donc encore ramené au cas précédent. D'ailleurs, nous supposons, dans ce qui va suivre, que la surface S ne présente pas de pareils points singuliers.

D'après les hypothèses faites, la surface S qui est algébrique n'a pas de nappes infinies; elle se compose donc d'un certain nombre de nappes fermées $S_1, S_2, \dots, S_p$, séparées les unes des autres. Au point de vue qui nous occupe, il nous suffira d'étudier séparément la forme des trajectoires sur une de ces nappes, sur la nappe S_1 par exemple.

Il est une notion qui va jouer un rôle fondamental dans ce qui va suivre, c'est le genre de la nappe S_1 au point de vue de la géométrie de situation.

Voici la définition de ce genre. Si l'on peut tracer, sur la surface fermée S_1 , p cycles fermés n'ayant aucun point commun, sans partager la surface en deux régions séparées, et si l'on n'en peut tracer davantage, on dira que la surface S_1 est de genre p (ou ce qui revient au même qu'elle est $2p + 1$ fois connexe). Ainsi la sphère est de genre 0, parce qu'on ne peut y tracer de cycle fermé sans partager sa surface en deux régions. Le tore est de genre 1, parce qu'un cercle méridien ou un cercle parallèle ne le divise pas en deux régions; et, si l'on a tracé sur la surface un cercle méridien, par exemple, on ne peut plus y tracer un nouveau cycle fermé, *ne rencontrant pas le premier*, sans partager le tore en deux régions.

Terminons ce Chapitre en étendant au cas qui nous occupe un théorème important de la première Partie.

Nous conserverons la convention faite au commencement de la deuxième Partie, c'est-à-dire que nous supposons que toute trajectoire qui va passer par un nœud est arrêtée à ce nœud, et que toute

trajectoire qui va passer par un col est continuée soit à droite, soit à gauche, par l'une des branches de courbe qui vont passer par ce col.

Cela posé, il est clair que les trajectoires peuvent se partager en quatre catégories :

- 1° Les cycles ou courbes fermées;
- 2° Les trajectoires qui sont arrêtées à un nœud;
- 3° Les trajectoires qui se terminent en tournant indéfiniment autour d'un foyer dont elles se rapprochent asymptotiquement;
- 4° Les trajectoires que l'on peut prolonger indéfiniment sans jamais revenir au point de départ, sans jamais rencontrer un nœud ou se rapprocher asymptotiquement d'un foyer.

Il est clair que la longueur de ces dernières est infinie, soit qu'on compte cette longueur d'arc sur la surface S_1 elle-même, soit qu'on la compte sur la projection de la courbe sur un plan quelconque.

Considérons maintenant une trajectoire qui ne rencontre aucun cycle algébrique qu'en un nombre fini de points. Il est évident qu'elle ne pourra appartenir à la troisième catégorie, car tout arc algébrique passant par un foyer rencontre en une infinité de points toute trajectoire qui tourne indéfiniment autour de ce foyer. Je dis qu'elle ne pourra pas non plus appartenir à la quatrième catégorie. Pour cela, je vais faire voir que, en supposant qu'une trajectoire de cette catégorie ne rencontre aucun cycle algébrique qu'en un nombre fini de points, on trouverait que la projection de cette trajectoire sur un certain plan devrait avoir une longueur finie, ce qui est contraire à ce que nous venons de voir.

En effet, considérons la portion de la trajectoire décrite par le point mobile depuis une certaine époque $t = t_0$ que nous déterminerons davantage plus loin, jusqu'à $t = +\infty$.

Nous partagerons la surface S_1 par un certain nombre de cycles algébriques en un certain nombre de régions, telles que chacune d'elles ne puisse être rencontrée qu'en un seul point par une γ parallèle à l'axe des z . Ces cycles algébriques ne seront rencontrés par la trajectoire qu'en un nombre fini de points. Donc nous pourrons prendre t_0 assez grand pour que, à partir de l'époque t_0 , le point mobile ne rencontre plus aucun de ces cycles et reste, par conséquent, à l'intérieur d'une des régions que nous venons de définir.

Il arrivera alors que, à partir de l'époque t_0 , la projection du point mobile sur le plan des xy restera constamment à l'intérieur d'une certaine région finie R de ce plan.

Considérons sur la surface S , le lieu des points, tels que la projection sur le plan des xy de la trajectoire qui passe par ce point présente un point d'inflexion. Ce lieu sera algébrique et, par conséquent, ne pourra être rencontré qu'en un nombre fini de points par la trajectoire que nous considérons. Nous pouvons donc prendre t_0 assez grand pour que, à partir de cette époque t_0 , la projection de cette trajectoire sur le plan des xy ne présente plus de point d'inflexion et soit, par conséquent, une courbe convexe.

Le lieu des points de la surface S , où l'on a $\frac{dx}{dt} = 0$, et celui des points où l'on a $\frac{dy}{dt} = 0$, sont encore algébriques. On peut en conclure, en raisonnant comme nous venons de le faire, que l'on peut prendre t_0 assez grand pour que, à partir de cette époque, $\frac{dx}{dt}$ et $\frac{dy}{dt}$ restent toujours de même signe, positifs par exemple.

Soient maintenant M_0 la projection du point mobile au temps t_0 , M_1 la projection de ce point au temps t_1 ($t_1 > t_0$). Par le point M_0 je mène une parallèle à l'axe des x ; par le point M_1 je mène une parallèle à l'axe des y rencontrant la première en P .

Soit R' un rectangle dont les côtés soient parallèles aux deux axes et qui soit tel que la région R définie plus haut y soit située tout entière. Le triangle curviligne M_0M_1P formé par l'arc de trajectoire M_0M_1 et les deux droites M_0P , M_1P sera *convexe* et situé tout entier à l'intérieur de R' . Son périmètre sera donc plus petit que celui de R' .

Donc l'arc M_0M_1 est toujours plus petit que le périmètre de R , et cela quel que soit le point M_1 . Donc la longueur de la projection de notre trajectoire serait finie, ce qui est absurde et nous oblige à rejeter l'hypothèse que la trajectoire soit de la quatrième catégorie.

D'où la conclusion suivante :

Toute trajectoire qui ne rencontre aucun cycle algébrique qu'en un nombre fini de points est un cycle fermé ou va aboutir à un nœud, où l'on doit l'arrêter.

CHAPITRE XIII.

DISTRIBUTION DES POINTS SINGULIERS.

Reprenons la nappe S_1 de genre p , et supposons que cette nappe ne présente ni point conique, ni courbe multiple.

Soient C le nombre des cols situés sur cette nappe, N le nombre des nœuds, F celui des foyers; je dis qu'on aura la relation

$$N + F - C = 2 - 2p.$$

Traçons sur la surface S_1 un cycle quelconque. Ce cycle sera touché en certains de ses points par diverses trajectoires, mais les unes le toucheront extérieurement, les autres intérieurement. Soient E le nombre des contacts extérieurs, I celui des contacts intérieurs; le nombre

$$J = \frac{E - I - 2}{2}$$

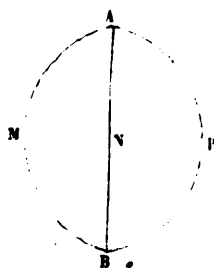
s'appellera l'*indice* du cycle. Si le cycle présente un point anguleux, il pourra arriver que la trajectoire qui passe par ce point traverse ce cycle en passant de l'extérieur à l'intérieur, auquel cas ce point ne doit pas compter pour un contact. Il pourra arriver aussi que cette trajectoire ne passe pas de l'extérieur du cycle à l'intérieur, mais reste constamment à l'extérieur si le point anguleux est saillant, ou constamment à l'intérieur si le point anguleux est rentrant. Alors le point anguleux devra compter pour un contact extérieur ou intérieur (*voir* la première Partie, p. 412). Nous supposons que le cycle a été choisi, de façon à partager la nappe S_1 en deux régions dont une au moins simplement connexe.

Si la nappe S_1 est de genre 0, les deux régions sont toutes deux simplement connexes, et il faudra une convention spéciale pour décider laquelle des deux doit être regardée comme l'intérieur du cycle.

Si la nappe S_1 est de genre > 0 , l'une des régions sera simplement connexe, et l'autre multiplément connexe, et ce sera la première que l'on considérera comme l'intérieur du cycle.

Cela posé, joignons deux points A et B par trois arcs de courbe AMB.

Fig. 1.



ANB, APB (fig. 1). Nous déterminerons ainsi trois cycles

$$C_1 = \text{ANBMA},$$

$$C_2 = \text{APBNA},$$

$$C_3 = \text{APBMA}.$$

Le troisième pourra être regardé comme la somme des deux autres

$$C_3 = C_1 + C_2.$$

Je dis que l'on aura

$$\text{ind. } C_3 = \text{ind. } C_1 + \text{ind. } C_2$$

ou, ce qui revient au même,

$$(1) \quad E_3 - I_3 = E_1 + I_1 - E_2 + I_2 = \dots = 2.$$

E_1, E_2, E_3 étant le nombre des contacts extérieurs, I_1, I_2, I_3 celui des contacts intérieurs des trajectoires avec les trois cycles.

Ces contacts se diviseront en cinq catégories :

- 1° Ceux qui ont lieu le long de AMB ;
- 2° Ceux qui ont lieu le long de APB.

Les premiers n'appartiennent qu'aux cycles C_1 et C_3 , les seconds n'appartiennent qu'aux cycles C_2 et C_3 . Un contact d'une de ces deux catégories entrera donc deux fois dans le premier membre de l'égalité (1), une fois avec le signe +, une fois avec le signe - ; les termes correspondants se détruiront.

- 3° Ceux qui ont lieu le long de ANB.

Un contact de cette catégorie est extérieur pour C_1 et intérieur pour C_2 , ou inversement. Il entrera donc deux fois dans le premier membre de l'égalité (1), une fois avec le signe $+$ et une fois avec le signe $-$. Les termes correspondants se détruiront.

4° Ceux qui peuvent avoir lieu en A.

Suivant la position de la trajectoire qui passe en A, on pourra avoir

Ou bien un contact extérieur pour les trois cycles;

Ou bien un contact extérieur pour le cycle C_1 seulement ou pour le cycle C_2 seulement;

Ou bien (dans le cas seulement où le point A serait un angle rentrant du cycle C_3) un contact intérieur pour le cycle C_3 ;

Ou bien encore (dans le cas seulement où le point A serait un angle rentrant des deux cycles C_1 et C_3) un contact intérieur pour les cycles C_1 et C_3 et un contact extérieur pour C_2 .

Dans tous les cas, la somme des termes correspondants du premier membre de (1) se réduira à -1 .

5° Les contacts qui ont lieu en B.

Pour la même raison, la somme des termes correspondants se réduira à -1 .

Le premier membre de (1) se réduit donc à -2 , de telle façon que cette égalité est vérifiée.

Donc l'indice d'un cycle total est égal à la somme des indices des cycles partiels qui le composent.

Cherchons maintenant à déterminer l'indice d'un cycle infiniment petit.

Si le cycle infiniment petit n'enveloppe aucun point singulier, nous pourrions toujours supposer qu'il est convexe, car, s'il ne l'était pas, on pourrait le décomposer en plusieurs cycles plus petits encore et convexes. Alors la *fig. 2* indique que le cycle a seulement deux contacts extérieurs avec les trajectoires Mp et $M''p''$.

L'indice est donc égal à 0.

Si le cycle enveloppe un col, nous le supposerons encore convexe, et la *fig. 3* montrera qu'il a quatre contacts extérieurs, et que son indice est égal à 1.

Si le cycle enveloppe un nœud ou un foyer, je dis que son indice est -1 . En effet, dans le voisinage d'un nœud ou d'un foyer, on

pourra toujours mener un cycle sans contact que nous supposons tout entier intérieur au cycle considéré. Ce cycle considéré pourra alors être décomposé en plusieurs autres, à savoir : le cycle sans contact dont il vient d'être question et d'autres cycles convexes n'enve-

Fig. 2.

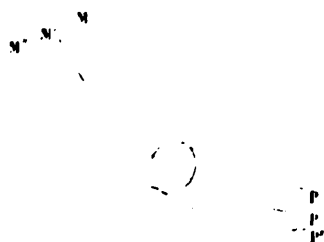


Fig. 3.



loppant pas le point singulier. L'indice du cycle sans contact sera égal à -1 ; celui des autres cycles, à 0 .

L'indice du cycle total sera donc -1 .

Il résulte de tout ce qui précède que l'indice d'un cycle quelconque est égal au nombre des cols qu'il contient diminué du nombre des nœuds et de celui des foyers.

Les centres qui sont des points singuliers exceptionnels rentreront, à ce point de vue, dans les foyers; en effet, autour d'un centre, on peut mener un cycle fermé avec deux contacts intérieurs et deux contacts extérieurs; ce cycle aura alors pour indice -1 .

Nous allons maintenant partager la surface de la nappe S_1 en un certain nombre de régions *simplement connexes* en y traçant un certain nombre de cycles.

La somme des indices de tous ces cycles sera évidemment

$$C - F - N.$$

Pour évaluer, d'une autre manière, cette somme d'indices, nous assimilerons à un polyèdre la figure formée par la nappe S_1 divisée en régions simplement connexes.

Tout le monde connaît le théorème d'Euler, d'après lequel, si α , β , γ sont le nombre des faces, des arêtes et des sommets d'un polyèdre

convexe, on doit avoir

$$\alpha = \xi + \gamma - 2.$$

Ce théorème s'étend aisément au cas où le polyèdre, au lieu d'être convexe, forme une surface de genre p ; on trouve alors

$$\alpha = \xi + \gamma - 2 + 2p.$$

Mais, en géométrie de situation, on n'a pas à s'inquiéter de la forme des faces et des arêtes; nous n'avons donc pas besoin de supposer que les faces du polyèdre sont planes, et ses arêtes rectilignes. Il en résulte que la figure, formée par la nappe S_i divisée en régions simplement connexes, est un véritable polyèdre curviligne auquel s'applique le théorème d'Euler. Les faces sont alors les régions simplement connexes elles-mêmes; une arête sera la portion du périmètre d'une de ces régions qui lui sert de frontière commune avec une région limitrophe; un sommet sera l'extrémité d'une arête, c'est-à-dire un point commun au périmètre de trois ou de plusieurs régions.

Supposons qu'un sommet soit commun au périmètre de v régions; il sera assimilable à un angle solide à v faces d'un polyèdre rectiligne. On aura, d'ailleurs,

$$\sum v = 2\xi.$$

Nous pourrions, d'ailleurs, supposer que les angles, formés par les diverses arêtes qui aboutissent à un sommet, sont tous saillants.

Nous cherchons à évaluer, pour l'ensemble de nos cycles, l'excès du nombre ΣE des contacts extérieurs sur le nombre ΣI des contacts intérieurs.

D'abord nous n'avons pas à nous préoccuper des contacts qui ont lieu en un point d'une arête; car, si un pareil contact est extérieur par rapport au cycle qui forme le périmètre d'une des deux régions auxquelles l'arête sert de frontière commune, il sera un contact intérieur pour le périmètre de la seconde de ces régions, et réciproquement.

Nous n'avons donc à considérer que les contacts qui peuvent avoir lieu aux sommets. Soit donc un sommet commun à v cycles. La trajectoire qui passe en ce point traversera deux de ces cycles et aura un

contact extérieur avec les c autres. On a donc

$$\Sigma E + \Sigma I = \Sigma c + 2 = \Sigma c + 2\gamma + 2\beta + \gamma.$$

La somme cherchée des indices est, d'ailleurs, égale à

$$\sum E + \frac{1}{2} \sum I = \frac{1}{2} (\sum E + \sum I) + \gamma$$

ou à

$$\beta + \gamma + \gamma = 2\gamma + 2.$$

Il vient donc

$$C + F + N = 2\gamma + 2. \quad \text{C. Q. F. D.}$$

Il résulte immédiatement de cette formule que les surfaces de genre 1 sont les seules qui puissent ne présenter aucun point singulier.

CHAPITRE XIV.

GÉNÉRALISATION DES DEUX PREMIÈRES PARTIES.

Nous allons reprendre maintenant chacun des théorèmes des Chapitres IV à VI pour voir s'ils s'étendent au cas qui nous occupe et dans quelle mesure ils doivent être modifiés.

Le théorème VI, d'après lequel tout cycle algébrique a un nombre de contacts pair, est encore vrai, mais seulement des cycles qui divisent la nappe S_1 en deux régions dont une au moins simplement connexe.

En effet, l'indice d'un pareil cycle, qui dépend du nombre des points singuliers qui y sont contenus, est essentiellement entier. Donc $\Sigma E = \Sigma I$ et, par conséquent, $\Sigma E + \Sigma I$, c'est-à-dire le nombre des contacts, sont toujours pairs.

Le théorème VII d'après lequel, si l'on peut mener entre deux points un arc *quelconque* sans contact, on peut aussi mener entre ces deux points un arc *algébrique* sans contact, est encore vrai pour les équations de degré supérieur. On n'a pour s'en assurer qu'à se reporter à la démonstration de la page 416, 1^{re} Partie. Nous aurons toutefois

une modification à y introduire; nous représenterons l'arc sans contact par les équations

$$x = \varphi(t), \quad y = \psi(t), \quad F(x, y, z) = 0,$$

les extrémités de cet arc correspondant à $t = 0$, $t = \pi$. La démonstration continuerait de la même façon que dans la 1^{re} Partie.

Il importe de remarquer que les séries

$$\begin{aligned} \frac{dx}{dt} &= \sum m A_m \cos mt - \frac{x_1}{2} \sin \frac{t}{2} + \frac{x_2}{2} \cos \frac{t}{2}, \\ \frac{dy}{dt} &= \sum m B_m \cos mt - \frac{y_1}{2} \sin \frac{t}{2} + \frac{y_2}{2} \cos \frac{t}{2} \end{aligned}$$

sont non seulement convergentes, mais uniformément convergentes, ce qui est nécessaire pour la suite de la démonstration; car la somme de la série

$$\sum m A_m \cos mt,$$

reprenant la même valeur pour t et pour $2\pi - t$, est une fonction continue de t quand cette variable croît de $-\infty$ à $+\infty$.

Le théorème VIII et sa démonstration subsistent aussi sans modification.

Si donc on peut joindre deux points A et B par un arc sans contact, et si A₁ et B₁ sont deux points des trajectoires qui passent par A et B, on pourra également joindre A₁ et B₁ par un arc sans contact.

Le théorème IX s'énonce ainsi :

Si AB et A₁B₁ sont deux arcs de trajectoires et si AA₁, BB₁ sont des arcs algébriques ne coupant AB et A₁B₁ en aucun autre point que A, B, A₁ et B₁, les nombres des contacts de AA₁ et de BB₁ seront de même parité.

Ce théorème subsistera encore dans le cas qui nous occupe, pourvu que le cycle ABA₁B₁ partage la nappe S₁ en deux régions, dont une simplement connexe.

THÉORÈME X. — *Si un arc de trajectoire qui ne passe par aucun point singulier est sous-tendu par un arc de courbe, le nombre des contacts de cet arc de courbe est impair.*

Ce théorème sera encore vrai, si le cycle formé par les deux arcs divise la nappe S_1 en deux régions, dont une simplement connexe.

Ainsi les théorèmes du Chap. IV, qui constituent ce que j'ai appelé la théorie des contacts, s'étendent avec quelques modifications au cas qui nous occupe. Je passe maintenant à la théorie des conséquents.

Soient

$$x = \varphi(t), \quad y = \psi(t)$$

un arc algébrique sans contact, et M_0, M_1 deux points consécutifs d'intersection de cet arc avec une même trajectoire. M_1 est le conséquent de M_0 , M_0 l'antécédent de M_1 , et si t_0 et t_1 sont les valeurs de t qui correspondent à ces deux points M_0 et M_1 , la relation qui lie t_1 et t_0 est la loi de conséquence.

Les théorèmes XI et XII ne subsistent qu'avec d'importantes modifications sur lesquelles nous reviendrons plus loin.

Le théorème XIII, au contraire, reste vrai pour les équations d'ordre supérieur.

Si $t_1 = \varphi_1(t_0)$ est la loi de conséquence, la fonction φ_1 est holomorphe. Il n'y a d'exception que pour les valeurs de t_0 qui correspondent à une trajectoire allant aboutir à un point singulier avant d'avoir rencontré de nouveau l'arc sans contact, et pour les valeurs de t_0 ou de t_1 qui correspondent aux extrémités de cet arc sans contact.

Le théorème XIV reste vrai également, mais la démonstration doit être modifiée, car le cycle $M_0 M_1 N_0 M_0$ dont il est question dans la démonstration donnée dans le Chap. V pourrait ne pas partager la nappe S_1 en deux régions. Mais soit P_0 un point situé sur l'arc sans contact entre N_0 et M_1 (voir *fig. 10*, p. 257, 11^e Partie) et à une distance finie de N_0 . Soit P_1 un point situé sur l'arc sans contact à droite de M_1 et à une distance finie de ce point. Joignons $P_0 P_1$ par un arc de courbe tel que le cycle fermé $M_0 M_1 P_1 P_0$ enferme une région simplement connexe. La trajectoire $N_0 N_1$ ne pourrait sortir de cette région sans s'éloigner de la trajectoire $M_0 M_1$ à une distance finie ou sans la traverser, ce qui est impossible. Donc le point N_1 doit être à droite du point M_1 .

C. Q. F. D.

Je dis qu'on peut toujours mener sur la nappe S_1 des cycles sans contact. En effet :

1° Dans le voisinage des nœuds et des foyers, on peut tracer des cycles sans contact enveloppant des points singuliers.

2° Si M_0 et M_1 sont deux points d'intersection consécutifs d'une trajectoire M_0PM_1 et d'un arc sans contact M_0QM_1 , le cycle

$$M_0QM_1PM_0$$

pourra être regardé comme sans contact. Soient en effet N_0 un point infiniment voisin de M_0 et à droite de ce point comme dans la *fig. 10* que nous venons de citer, N_1 le conséquent du point N_0 . Nous pourrions tracer dans la région infiniment mince comprise entre les deux trajectoires M_0M_1 et N_0N_1 un arc N_0RM_1 ne coupant chaque trajectoire qu'une fois. Le cycle $N_0RM_1QN_0$, qui diffère infiniment peu de $M_0PM_1QM_0$, sera alors sans contact.

Ainsi donc, si la nappe S_1 contient un nœud ou un foyer, on sera certain de pouvoir tracer un cycle sans contact; si elle n'en contient pas, toute trajectoire devra couper au moins un cycle algébrique en une infinité de points (à moins de se réduire à un cycle fermé, comme nous l'avons vu dans le Chap. XII). Ce cycle algébrique pourra être partagé en un nombre fini d'arcs sans contact. Donc au moins un de ces arcs sans contact sera coupé par la trajectoire en une infinité de points. Parmi ces points, nous retiendrons deux points consécutifs M_0 et M_1 et ils nous donneront un cycle sans contact, ainsi qu'on vient de le voir. *Il y a donc toujours des cycles sans contacts à moins que toutes les trajectoires ne se réduisent à des cycles fermés.*

Si un cycle sans contact divise la nappe S_1 en deux régions dont une simplement connexe, cette dernière contient au moins un nœud ou un foyer.

En effet, l'indice du cycle sans contact est égal à -1 . Si donc N , F et C désignent le nombre des nœuds, des foyers et des cols contenus à l'intérieur de notre région simplement connexe, on aura

$$N + F - C = 1,$$

ce qui démontre le théorème énoncé.

Supposons maintenant que le cycle sans contact divise la nappe S , en deux régions pouvant toutes deux être multiplement connexes. Je dis qu'il y aura des points singuliers dans chacune de ces deux régions.

Soit R l'une de ces régions que je supposerai q fois connexe. Il s'agit de trouver l'indice du cycle sans contact C en considérant la région R comme l'intérieur de ce cycle. Nous pourrions construire une calotte R_1 simplement connexe et admettant pour frontière le cycle C . L'ensemble des deux régions R et R_1 formera alors une surface fermée de genre $\frac{q-1}{2}$, ce qui prouve en passant que q doit être impair (cf. RIEMANN, *Gesammelte Werke*, p. 12; Leipzig, Teubner, 1876). Posons donc

$$q = 2h + 1.$$

Si nous décomposons la région R en un certain nombre de régions simplement connexes, la surface fermée $R + R_1$ pourra être assimilée à un polyèdre. En appliquant à ce polyèdre le théorème d'Euler et en raisonnant comme dans le Chapitre précédent, on trouvera

$$\text{indice de } C = \Sigma i - 2h,$$

Σi désignant la somme des indices des cycles à l'aide desquels la région R a été partagée en régions simplement connexes. Or, si N , F et C sont les nombres des nœuds, des foyers et des cols de la région R , on a

$$\Sigma i = C - N - F.$$

De plus, le cycle C étant sans contact, il vient

$$\text{indice de } C = -1;$$

d'où

$$C - N - F = 2h - 1$$

ou

$$N + F + C \equiv 1, \quad \text{mod. } 2,$$

ce qui prouve que N , F et C ne peuvent pas être nuls à la fois.

Au point de vue qui nous occupe en ce moment, toute trajectoire fermée et ne passant par aucun point singulier pourra être assimilée

à un cycle sans contact, je veux dire que son indice sera égal à -1 . Une trajectoire fermée n'a pas d'indice à proprement parler, mais les cycles qui en diffèrent infiniment peu en auront un, et je dis que cet indice sera -1 .

En effet, il peut se présenter deux cas :

Soit M_0PM_0 une trajectoire fermée, soit AM_0B un arc sans contact, soit t le paramètre qui définit la position d'un point sur cet arc; soit 0 la valeur de t qui correspond à M_0 . Soit

$$t_1 = \varphi_1(t_0)$$

la loi de conséquence sur notre arc. On aura

$$\varphi_1(0) = 0.$$

Il pourra se faire d'abord que la fonction φ_1 ne soit pas identiquement nulle. Dans ce cas, soient t_0 une valeur infiniment petite de t , N_0 le point correspondant, N_1 son conséquent et t_1 la valeur infiniment petite correspondante de t . Soit N_0QN_1 l'arc de trajectoire qui joint N_0 à N_1 . Le cycle $N_0QN_1N_0$ qui différera infiniment peu de M_0PM_0 pourra être assimilé, d'après ce qu'on a vu plus haut, à un cycle sans contact. Dans ce cas, la trajectoire fermée M_0PM_0 est un *cycle limite* et jouit des propriétés de ces cycles démontrées dans la II^e Partie.

Il peut arriver aussi que la fonction φ_1 soit identiquement nulle. Dans ce cas les trajectoires voisines de M_0PM_0 sont des cycles fermés s'enveloppant mutuellement.

Si alors on mène un cycle infiniment peu différent de M_0PM_0 le nombre de ses contacts extérieurs sera le même que celui de ses contacts intérieurs et son indice sera encore -1 . C. Q. F. D.

Si donc une trajectoire fermée partage la nappe S_1 en deux régions, il y aura des points singuliers dans chacune d'elles.

Donc tout cycle algébrique qui passe par tous les points singuliers rencontre tous ceux des cycles sans contact et toutes celles des trajectoires fermées qui partagent la nappe S_1 en deux régions.

C'est la généralisation du théorème XVI de la deuxième Partie.

Passons maintenant à la généralisation du théorème XVIII qui est l'objet principal de ce Chapitre. Ce théorème généralisé s'énonce ainsi :

On peut sillonner la nappe S_1 par une série de cycles et de polycycles courbes fermées à point double, de telle façon que par chaque point de cette nappe passe un de ces cycles et un seul, excepté par les foyers et les nœuds qui seront regardés comme des cycles infiniment petits. Parmi ces cycles, les uns seront des cycles sans contact, les autres des trajectoires fermées.

Si le genre p de la nappe S_1 est égal à 0, la généralisation est immédiate. En effet, pour démontrer sur la sphère les théorèmes X à XVIII des deux premières Parties, nous nous sommes appuyés seulement sur une propriété de la sphère, celle d'être simplement connexe ou de genre 0.

Il est encore un autre cas où cette généralisation peut se faire immédiatement. Supposons qu'on ait découpé sur la nappe S_1 une région R doublement connexe et limitée par deux cycles sans contact C et C' . Supposons de plus que cette région R ne renferme aucun point singulier. Je dis que le théorème XVIII s'appliquera à cette région, c'est-à-dire qu'on pourra sillonner cette région par une infinité de cycles K qui seront tous des cycles sans contact ou des trajectoires fermées.

Il est clair d'ailleurs que ces cycles K devront être disposés de façon à partager la région R en deux autres, la première, limitée par les cycles K et C , la seconde par les cycles K et C' . En d'autres termes, il n'est pas possible que le cycle K forme à lui seul la frontière d'une région R' contenue tout entière dans R . En effet, l'indice de ce cycle est égal à -1 , tandis que R , et par conséquent R' , ne renfermerait pas de point singulier.

Divisons maintenant, à l'aide d'un point M quelconque, toute trajectoire en deux demi-trajectoires analogues aux demi-caractéristiques des deux premières parties. L'une de ces demi-trajectoires comprendra les points où le point mobile passera après être passé au point M et l'autre les points où le point mobile était passé avant d'arriver au point M .

Les demi-trajectoires se diviseront alors en trois catégories :

1° Les trajectoires fermées qui resteront tout entières à l'intérieur de R ;

2° Les demi-trajectoires qui s'étendront indéfiniment sans jamais se fermer, et sans jamais sortir de R;

3° Les demi-trajectoires qui sortent de R par l'un des cycles C et C'.

Nous commencerons par entourer les trajectoires fermées dans des *anneaux limites* ainsi que nous l'avons fait dans la II^e Partie (p. 270). A cet effet, joignons un point de C à un point de C' par un arc algébrique quelconque. Cet arc algébrique pourra être décomposé en un nombre fini d'arcs sans contact. Par chacune des extrémités de ces divers arcs je fais passer un petit arc sans contact. J'ai ainsi tracé à l'intérieur de R un certain nombre d'arcs sans contact et je suis certain que toute demi-trajectoire de la première catégorie rencontrera au moins un d'entre eux.

Considérons un quelconque de ces arcs sans contact et cherchons quels sont ceux des points de cet arc qui ont un conséquent. Nous convenons pour cela que, si la trajectoire qui passe par un point M_0 de l'arc sans contact sort de la région R ou si elle ne vient plus couper de nouveau l'arc sans contact, le point M_0 sera regardé comme sans conséquent. Si, au contraire, la trajectoire qui passe par M_0 vient couper de nouveau l'arc sans contact en un point M_1 , sans être sortie de R, le point M_1 sera le conséquent de M_0 .

Parmi les arcs sans contact, en nombre fini, que j'ai tracés dans la région R, les uns ne seront rencontrés par aucune trajectoire fermée de la première catégorie et devront être rejetés, les autres seront rencontrés par au moins une trajectoire fermée; je les appellerai *arcs auxiliaires*.

Considérons un arc auxiliaire quelconque; sur cet arc on pourra toujours trouver des points admettant un conséquent; car le point où il est coupé par une trajectoire fermée est son propre conséquent. De plus aucune trajectoire issue d'un point de l'arc auxiliaire n'ira passer par un col, puisque la région R n'en renferme pas. La courbe de conséquence affectera donc l'une des formes indiquées sur la *fig. 11* (II^e Partie, p. 258).

Nous avons vu que, si l'on joint un point M_0 d'un arc sans contact à son conséquent M_1 par un arc de trajectoire M_0PM_1 , le cycle

$$M_0PM_1M_0$$

peut être regardé comme un cycle sans contact. Il y aurait exception, bien entendu, si les deux points M_0 et M_1 se confondaient, auquel cas le cycle sans contact se réduirait à une trajectoire fermée. Nous traçons les cycles sans contact ainsi obtenus pour tous les points de l'arc auxiliaire qui admettent un conséquent. L'ensemble de ces cycles formera alors un anneau limite comme ceux que nous avons envisagés dans la II^e Partie.

Cet anneau limite sera sillonné par une série de cycles sans contact dont quelques-uns se réduiront à des trajectoires fermées. De plus cet anneau limite partagera la région R en deux autres régions analogues R' et R'' , ne renfermant plus qu'un nombre moindre d'arcs auxiliaires.

En continuant de la sorte sur les deux régions R' et R'' , on arrivera à partager la région R en un certain nombre d'anneaux limites et en un certain nombre de régions interannulaires à l'intérieur desquelles on ne pourra plus tracer aucune trajectoire fermée (cf. II^e Partie, p. 272).

Considérons une de ces régions interannulaires; elle sera tout à fait analogue à la région R ; seulement on n'y pourra pas tracer de demi-trajectoires de la première catégorie. Je dis qu'on n'en pourra pas non plus tracer de la deuxième.

En effet, toute demi-trajectoire de la deuxième catégorie admet un cycle limite qui ne pourrait être qu'une demi-trajectoire de la première; car, dans l'intérieur de la région R , les théorèmes du Chapitre V s'appliquent sans restriction; donc une demi-trajectoire ne peut rester constamment à l'intérieur de la région interannulaire qui n'est traversée par aucune trajectoire de la première catégorie.

Ainsi une région interannulaire ne peut renfermer que des demi-trajectoires de la troisième catégorie; d'où il suit que toute trajectoire qui traverse cette région va aboutir de part et d'autre à deux extrémités situées sur les deux cycles γ et γ' qui limitent la région. Il n'est pas possible que les deux extrémités soient sur un même cycle γ ; car un arc sans contact ne peut sous-tendre un arc de trajectoire.

Donc toute trajectoire tracée dans la région interannulaire va d'un point du cycle γ à un point du cycle γ' ; ce qui démontre la possibilité de sillonner cette région de cycles sans contact.

Le théorème XVIII est donc démontré pour la région R .

Un cas particulier digne d'intérêt est celui où la loi de conséquence

sur un des arcs auxiliaires s'écrit $t_1 = t_0$. On reconnaîtrait alors par un raisonnement analogue à celui que nous avons employé à la fin du Chap. XI en remarquant que la région R ne renferme aucun point singulier et que toutes les trajectoires qui traversent cette région sont fermées.

Si on laisse de côté ce cas particulier, toutes les trajectoires fermées sont des cycles limites; et le théorème XVII, d'après lequel ces cycles limites sont en nombre fini, se généralise aisément. (Il ne s'agit jusqu'ici, bien entendu, que des cycles limites et trajectoires fermées situés tout entiers à l'intérieur de R.)

Les théorèmes XVII et XVIII sont encore vrais quand un des cycles C et C' qui limitent la région R, au lieu d'être un cycle sans contact, devient une trajectoire fermée. Il y aurait exception toutefois pour le théorème XVII si le cycle C se réduisait à une trajectoire fermée allant passer par un col.

Il me reste, pour démontrer les théorèmes XVII et XVIII dans toute leur généralité, à faire voir la possibilité de décomposer la nappe S_1 en un certain nombre de régions telles que R.

Pour cela, nous allons d'abord faire quelques remarques préliminaires.

Nous avons vu que si l'on joint deux points M_0 et M_1 d'un arc sans contact par un arc de trajectoire $M_0 P M_1$, le cycle $M_0 M_1 P M_0$ peut être regardé comme sans contact, c'est-à-dire que par chacun de ses points on peut mener un cycle sans contact qui diffère infiniment peu du cycle $M_0 M_1 P M_0$.

De même, supposons qu'un arc de trajectoire $M_0 Q M_1$ vienne aboutir en M_1 à un cycle sans contact $M_1 P M_1$. Je dis que le cycle

$$M_1 P M_1 Q M_0 Q M_1$$

pourra être regardé comme sans contact.

En effet, soient $M_0' Q' M_1'$, $M_0'' Q'' M_1''$, ..., $M_0^k Q^k M_1^k$ des arcs de trajectoire infiniment peu différents de $M_0 Q M_1$. Nous pourrions toujours tracer un arc de courbe, quittant le cycle sans contact $M_1 P M_1$ en M_1' , coupant chacun de ces arcs de trajectoire en un seul point, allant passer par le point M_0 et venant rejoindre le cycle sans contact en M_1^k .

Le cycle fermé $M_1^* P M_1^* Q^* M_0 M_1^*$, qui diffère infiniment peu de

$$M_1 P M_1 Q M_0 Q M_1,$$

sera alors sans contact.

Imaginons maintenant qu'on ait tracé sur la nappe S_1 un certain nombre de cycles sans contact (ou de trajectoires fermées), et qu'on ait déterminé ainsi sur cette nappe une certaine région P limitée par ces divers cycles sans contact. (Il peut se faire que la région P contienne la nappe S_1 tout entière; ainsi supposons que S_1 soit un tore, et qu'on ait tracé sur ce tore un cercle méridien C qui soit un cycle sans contact. La nappe S_1 forme alors tout entière une région P limitée par *deux* cycles, car on devra distinguer les deux lèvres de la coupure que l'on a faite sur le tore.

Cela posé, je dis que par tout point M_0 de la région P on peut tracer à l'intérieur de cette région un cycle sans contact. Considérons, en effet, la demi-trajectoire qui passe par M_0 . Voici les cas qui pourront se présenter :

1° Il pourra arriver que la demi-trajectoire se ferme sans être sortie de la région P . C'est le seul cas d'exception; on ne pourra pas faire passer par le point M_0 de cycle sans contact, mais seulement une trajectoire fermée.

2° Ou bien que la demi-trajectoire $M_0 Q M_1$ sorte de la région P en M_1 par un des cycles sans contact qui la limitent, par exemple par le cycle $M_1 H M_1$. Dans ce cas, le cycle $M_1 H M_1 Q M_0 Q M_1$ peut être regardé comme sans contact, ainsi qu'on vient de le voir.

3° Ou bien que la demi-trajectoire aboutisse à un nœud ou à un foyer. Ce cas se ramène au précédent, car on peut entourer le nœud ou le foyer d'un petit cycle sans contact, que la trajectoire est obligée de traverser pour aboutir au nœud ou au foyer.

4° Ou bien que la demi-trajectoire s'étende indéfiniment, sans se fermer, sans aboutir à un nœud ou à un foyer, et sans sortir de la région P . On pourra alors trouver, dans la région P , un arc sans contact qui coupe la demi-trajectoire en deux points N_0 et N_1 . Il s'ensuit, d'après le théorème VIII, qu'on pourra joindre M_0 à un autre point M_1 de la demi-trajectoire par un arc sans contact. Dans ce cas, le cycle

formé par l'arc sans contact M_0M_1 et par l'arc de trajectoire M_0M_1 , pourra être regardé comme sans contact.

Ainsi l'on pourra toujours mener, par le point M_0 et dans la région P , un cycle sans contact qui pourra, dans certains cas, se réduire à une trajectoire fermée.

Si le point M_0 est un col, il aboutit à ce col, non par deux, mais par quatre demi-trajectoires. On peut alors mener par le point M_0 deux cycles sans contact formant, par leur ensemble, une courbe fermée à point double.

Cela posé, voici comment nous opérerons pour partager la nappe S_1 en régions analogues à R . Nous commencerons par envelopper tous les nœuds et tous les foyers par des cycles infiniment petits sans contact. Soit dans la région P ainsi déterminée un col quelconque; par ce col, je pourrai faire passer deux cycles sans contact; j'aurai alors divisé la nappe S_1 en régions P plus petites; j'opérerai de même sur chacun des cols, et j'aurai finalement partagé la nappe S_1 en un certain nombre de régions R limitées par des cycles sans contact et ne contenant aucun point singulier.

Je dis que chacune de ces régions est doublement connexe et limitée par deux cycles seulement.

Supposons, en effet, que la région R soit q fois connexe et limitée par n cycles.

Nous aurons d'abord (voir RIEMANN, *loc. cit.*, p. 12).

$$q > n > 0, \quad q \equiv n \pmod{2}.$$

De plus, chacun des cycles a pour indice -1 , et il n'y a pas de point singulier dans la région. Si nous fermons la région R en construisant, sur chacun des cycles C qui la limitent, une calotte simplement connexe, puis que nous divisons la région elle-même en domaines simplement connexes par des cycles C' , nous obtiendrons une sorte de polyèdre curviligne qui sera de genre $\frac{q-n}{2}$. Le théorème d'Euler nous donnera donc

$$\Sigma i + \Sigma i' = q - n - 2,$$

si l'on désigne par Σi la somme des indices des cycles C et par $\Sigma i'$

la somme des indices des cycles C' . Or on a

$$\Sigma i = -n, \quad \Sigma i' = 0,$$

d'où

$$q = 2, \quad n = 2. \qquad \text{C. Q. F. D.}$$

Ainsi toutes nos régions sont doublement connexes et limitées par deux cycles : nous pouvons donc sillonner chacune d'elles de cycles sans contact, ce qui démontre le théorème XVIII dans toute sa généralité.

Quand on aura construit sur la nappe S_1 une série de cycles sans contact et de cycles limites s'enveloppant mutuellement, on pourra s'en servir pour construire les trajectoires elles-mêmes.

Si un cycle sans contact divise S_1 en deux régions, il ne pourra être coupé en plus d'un point par une même trajectoire.

Cela ne sera plus vrai, au contraire, si le cycle sans contact ne partage pas S_1 en deux régions; mais, même dans ce cas, la considération des cycles sans contact n'en conserve pas moins une importance capitale. C'est ce que l'on comprendra mieux par les exemples que je vais donner dans le Chapitre suivant.

CHAPITRE XV.

ÉTUDE PARTICULIÈRE DU TORE.

Les surfaces de genre 1 sont, comme on l'a vu, les seules sur lesquelles il puisse n'exister aucun point singulier. Après le cas des surfaces de genre zéro, qui ne diffère pas en réalité de ceux qui ont été traités dans les deux premières parties, le cas le plus simple qui puisse se présenter est celui des surfaces de genre 1 sans point singulier. C'est donc celui que nous allons étudier en détail.

Pour fixer davantage encore les idées, je supposerai que la surface S_1 se réduit au tore

$$z^2 + (R - \sqrt{x^2 + y^2})^2 = r^2.$$

Mais tout ce que nous dirons du tore s'appliquera à une surface

quelconque du genre 1, qui n'en diffère pas au point de vue de la géométrie de situation.

Nous poserons

$$\begin{aligned}x &= (R + r \cos \omega) \cos \varphi, \\y &= (R + r \cos \omega) \sin \varphi, \\z &= r \sin \omega.\end{aligned}$$

Nous mettrons les équations différentielles sous la forme

$$\frac{d\omega}{dt} = \Omega, \quad \frac{dz}{dt} = \Phi.$$

Ω et Φ devront être des polynômes entiers en $\cos \omega$, $\sin \omega$, $\cos \varphi$ et $\sin \varphi$.

On trouve, d'ailleurs,

$$\begin{aligned}\cos \varphi &= \frac{x(x^2 + y^2 + z^2 + R^2 - r^2)}{2R(x^2 + y^2)}, \quad \sin \varphi = \cos \varphi \frac{y}{x}, \\ \sin \omega &= \frac{z}{r}, \quad \cos \omega = \frac{x \cos \varphi + y \sin \varphi - R}{r},\end{aligned}$$

ce qui montre que si Ω et Φ sont des polynômes entiers en $\cos \omega$, $\sin \omega$, $\cos \varphi$ et $\sin \varphi$, $\frac{dx}{dt}$, $\frac{dy}{dt}$ et $\frac{dz}{dt}$ seront des fonctions rationnelles en x , y et z .

Pour bien saisir la suite de la discussion, il faut se reporter aux Chapitres VIII et IX (II^e Partie). Nous allons, en effet, appliquer les mêmes méthodes et partager le tore en régions acycliques (sillonnées par des cycles sans contact, parmi lesquels il n'y a aucun cycle limite), en régions cycliques (sillonnées par des cycles sans contact, parmi lesquels il y a certainement des cycles limites), en régions monocycliques (où il y a un cycle limite *et un seul*), et en régions douteuses (où l'on ne saurait affirmer qu'il y ait ni qu'il n'y ait pas de cycle limite). La discussion sera terminée lorsqu'il ne restera plus que des régions acycliques et monocycliques.

Le problème revient donc à savoir si une région est cyclique, et si elle est monocyclique. Voici les règles que nous appliquerons pour le

reconnaitre, et qui ne seront autre chose que celles que nous avons exposées dans le Chapitre VIII.

1^o Soit une région R doublement connexe limitée par deux cycles sans contact C et C' , et soit AMB un arc algébrique quelconque coupant ces deux cycles en A et en B . Considérons un observateur suivant cet arc et pénétrant dans la région R par le point A . S'il a à sa droite la demi-trajectoire qui s'étend dans la région R à partir du point A , nous dirons que le cycle C est positif, il sera négatif dans le cas contraire. Supposons maintenant que l'observateur, en suivant cet arc, sorte de la région R par le point B . De ce point B partiront deux demi-trajectoires, l'une s'étendant à l'intérieur de R , l'autre à l'extérieur de cette région. C'est cette dernière que nous considérerons. Si elle est à la droite de l'observateur, le cycle C' sera positif (*cf.* Chap. VIII, p. 286).

Si les deux cycles sont de même signe, le nombre des cycles limites contenus dans la région R est de même parité que le nombre des contacts de l'arc AMB ; il est de parité différente dans le cas contraire.

En particulier, supposons que les deux cycles C et C' soient deux cercles méridiens $z = z_0$ et $z = z_1$. Supposons que, dans la région R et sur sa frontière, Ω soit constamment de même signe; supposons que le long du cycle C on ait

$$\frac{dz}{dt} > 0$$

et que le long du cycle C' on ait

$$\frac{dz}{dt} < 0.$$

Prenons pour l'arc AMB l'arc de parallèle $\omega = 0$, qui est sans contact. Les deux cycles seront de signe contraire. Le nombre des cycles limites contenus dans R sera donc impair. Donc cette région est certainement cyclique.

2^o Soient ρ et θ les coordonnées polaires d'un point dans un plan; supposons qu'on établisse entre ω et z d'une part, ρ et θ d'autre part, une relation telle que la région R du tore et une certaine région R_1 doublement connexe du plan se correspondent point par point et

d'une façon uniforme. Supposons que, toutes transformations faites, l'équation différentielle proposée s'écrive

$$\frac{dz}{d\theta} = \varphi(\rho, \theta).$$

Si dans la région R il y a plus de deux cycles limites, il y aura forcément dans la région des points où

$$\frac{dz}{d\varphi} = 0 \quad \text{ou bien} \quad \varphi = \infty$$

(cf. Chap. XIII, théorème XIX).

Soient, en particulier,

$$\zeta = \omega, \quad \varphi = \varphi + K,$$

K étant une constante quelconque.

Si dans la région R la fonction Ω ne s'annule pas, ni non plus la fonction $\frac{d\Omega}{d\varphi}$, la région R est certainement acyclique ou monocyclique.

Nous étudierons d'abord l'exemple suivant

$$\frac{dz}{dt} = a + b \cos \omega - \cos \varphi, \quad \frac{d\omega}{dt} = c,$$

a, b, c étant des constantes telles que

$$a > b > 0, \quad a + b < 1, \quad c > 0.$$

Nous poserons

$$a + b = \cos \varphi_0, \quad a - b = \cos \varphi_1,$$

les angles φ_0 et φ_1 étant compris entre 0 et $\frac{\pi}{2}$.

La valeur de $\frac{d\omega}{dt}$ ne s'annulant jamais, tous les parallèles $\omega = \text{const.}$ sont des cycles sans contact. Voilà donc un premier système de cycles sans contact parmi lesquels il n'y a aucun cycle limite.

Maintenant, on voit aisément que, si φ est compris entre φ_1 et $2\pi - \varphi_1$,

$\frac{dz}{dt}$ est positif, et que, si φ est compris entre $-\varphi_0$ et $+\varphi_0$, $\frac{dz}{dt}$ est négatif.

Le tore se trouve ainsi partagé en quatre régions, les régions

$$\varphi_1 < \varphi < 2\pi - \varphi_1, \quad -\varphi_0 < \varphi < \varphi_0,$$

où tous les méridiens sont des cycles sans contact (ce sont des régions acycliques), et les régions $R(\varphi_0 < \varphi < \varphi_1)$ et $R'(-\varphi_1 < \varphi < -\varphi_0)$ qui sont douteuses.

Considérons, en particulier, la région R . Je dis d'abord qu'elle est cyclique, car on a

$$\frac{dz}{dt} > 0 \quad \text{pour le cycle } \varphi = \varphi_1,$$

$$\frac{dz}{dt} < 0 \quad \text{pour le cycle } \varphi = \varphi_0,$$

et, de plus,

$$\frac{d\omega}{dt} > 0.$$

Les deux cycles sans contact $\varphi = \varphi_1$ et $\varphi = \varphi_0$ qui limitent la région R sont donc de signe contraire, si l'on prend pour l'arc AMB , dont il a été question plus haut, l'arc sans contact $\omega = 0$. Donc la région R contient au moins un cycle limite.

De plus, elle n'en peut contenir plus d'un; car, si elle en contenait deux, on devrait avoir à l'intérieur de R soit $\Omega = 0$, soit $\frac{d\Phi}{d\varphi} = 0$.

Or Ω ne s'annule jamais, et $\frac{d\Phi}{d\varphi} = \sin \varphi$ ne peut s'annuler que pour $\varphi = m\pi$, c'est-à-dire en dehors de R .

Donc dans la région R on peut tracer un cycle limite et un seul, que j'appelle C .

Pour la même raison, dans la région R' , on peut tracer un cycle limite et un seul que j'appelle C' .

Nous sommes donc conduits à un second système de cycles sans contact parmi lesquels il y a deux cycles limites C et C' et une infinité de méridiens.

Les deux cycles C et C' partagent le tore en deux régions P et P' , toutes deux sillonnées de cycles sans contact. Considérons le point

mobile dont le mouvement est défini par nos équations différentielles. Si, à l'origine du temps, ce point mobile est dans une des deux régions P ou P' , il n'en pourra jamais sortir. Si donc sa coordonnée φ est à l'origine du temps comprise entre φ_1 et $2\pi - \varphi_1$, elle restera toujours comprise entre φ_0 et $2\pi - \varphi_0$.

De plus, dès que le point mobile aura franchi un des cycles sans contact du second système, il ne pourra plus le franchir de nouveau. Quand le temps croîtra indéfiniment, le point mobile se rapprochera asymptotiquement de l'un des deux cycles limites C et C' .

En d'autres termes, et pour reprendre le langage du Chapitre X, l'orbite du point mobile sera instable.

Le second exemple que nous traiterons sera encore plus simple que le précédent.

J'écrirai simplement

$$\frac{d\omega}{dt} = a, \quad \frac{dz}{dt} = b,$$

a et b étant deux constantes positives. On voit immédiatement que tous les parallèles et tous les méridiens sont des cycles sans contact, et que l'intégrale générale s'écrit

$$\frac{\omega}{a} = \frac{z}{b} = \text{const.}$$

Si le rapport $\frac{a}{b}$ est commensurable, toutes les trajectoires sont fermées; il y a donc stabilité.

Supposons maintenant que le rapport $\frac{a}{b}$ soit incommensurable; il y aura encore stabilité en ce sens, que si l'on considère une portion p du tore, si petite qu'elle soit, cette portion sera traversée une infinité de fois par une quelconque des trajectoires. Si le point mobile occupe à l'origine des temps un certain point A , il ne viendra pas repasser en ce point, mais il viendra une infinité de fois passer en des points infiniment voisins.

Considérons la courbe

$$\omega = c\varphi + d,$$

c étant une constante commensurable et d une constante quelconque.

Cette courbe sera un cycle fermé, et il est aisé de voir que ce sera toujours un cycle sans contact.

On peut donc tracer sur le tore une infinité de systèmes de cycles sans contact, parmi lesquels il n'y a aucun cycle limite.

Les deux exemples qui précèdent suffisent déjà pour faire comprendre que la présence d'un cycle limite est un signe d'instabilité, et l'absence d'un pareil cycle, un signe de stabilité. Mais, pour mieux se rendre compte de ce fait important, il est indispensable de pénétrer plus avant dans la question.

Nous supposons, dans ce qui va suivre, que ω et φ vont constamment en croissant avec le temps, c'est-à-dire que Ω et Φ sont toujours positifs, et de plus qu'il n'y a sur le tore aucun point singulier.

Considérons le méridien $\varphi = 0$, qui sera un cycle sans contact. Soit $M(0)$ un point de ce méridien où se trouve le point mobile à l'origine des temps. Ce point aura pour coordonnées

$$\varphi = 0, \quad \omega = \omega_0.$$

Si l'on fait croître le temps, ω et φ croîtront également, de sorte que φ finira par devenir égal à 2π : le point mobile sera venu alors en un point $M(1)$ qui sera situé sur le cycle sans contact $\varphi = 0$ d'où nous sommes partis, qui sera le conséquent du point $M(0)$ et qui aura pour coordonnées

$$\varphi = 2\pi, \quad \omega = \omega_1.$$

Les deux quantités ω_0 et ω_1 seront liées par une certaine relation qui n'est autre chose que la loi de conséquence du Chapitre V. J'écrirai cette loi sous la double forme

$$\omega_1 = \psi(\omega_0), \quad \omega_0 = \psi(\omega_1).$$

Les hypothèses faites permettent d'énoncer au sujet de cette loi de conséquence les principes suivants :

Premier principe. — La fonction ψ est continue.

Deuxième principe. — La fonction ψ croît constamment avec ω_0 , de telle sorte que

$$\frac{d\omega_1}{d\omega_0} > 0,$$

et, de plus, on a

$$\psi(\omega_0 + 2\pi) = \psi(\omega_0) + 2\pi.$$

Troisième principe. — La fonction ψ est holomorphe pour toutes les valeurs réelles de ω_0 .

D'ailleurs il est clair que la fonction ζ jouit des mêmes propriétés que la fonction ψ .

Soient maintenant $M(1)$, $M(2)$, ..., $M(i)$ les conséquents successifs et $M(-1)$, $M(-2)$, ..., $M(-i)$ les antécédents successifs de $M(0)$. Leurs coordonnées ω_1 , ω_2 , ... et ω_{-1} , ω_{-2} , ... nous seront données par les équations

$$\begin{aligned}\omega_1 &= \psi(\omega_0), & \omega_2 &= \psi\psi(\omega_0), & \omega_3 &= \psi\psi\psi(\omega_0), & \dots \\ \omega_{-1} &= \theta(\omega_0), & \omega_{-2} &= \theta\theta(\omega_0), & \omega_{-3} &= \theta\theta\theta(\omega_0), & \dots\end{aligned}$$

Les points M forment un ensemble de points que j'appellerai P , d'après la notation adoptée par M. Cantor. J'appellerai P' l'ensemble dérivé de P , c'est-à-dire l'ensemble des points dans le voisinage desquels il y a une infinité de points appartenant à l'ensemble P . Enfin je désignerai par

$$D(P, Q)$$

l'ensemble des points communs à deux ensembles P et Q . Je renverrai, pour plus de détails sur les ensembles de points et l'emploi des notations qui précèdent, à divers Mémoires de M. Cantor, parus en allemand dans les *Mathematische Annalen* et en français dans les *Acta mathematica*.

Je puis alors me poser les questions suivantes :

1° Quelles sont les propriétés de l'ensemble P des points M et de ses dérivés?

2° Dans quel ordre circulaire sont distribués les points M sur le cercle méridien $\varphi = 0$?

Je vais démontrer d'abord que l'on a

$$D(P, P') = 0$$

ou bien

$$D(P, P') = P.$$

En d'autres termes, si l'ensemble P a un point commun avec son dérivé P' , tous ses points feront partie de cet ensemble dérivé.

En effet, soit $M(i)$ un point de P appartenant à P' ; d'après la définition de P' , il y aura, dans le voisinage de $M(i)$, une infinité de points appartenant à P . Nous pourrions donc trouver une série de points

$$M(k_1), M(k_2), \dots, M(k_p), \dots$$

appartenant à P , et tels que

$$\lim M(k_p) = M(i) \quad \text{pour } p = \infty.$$

Cela posé, soit $M(v)$ un point quelconque de P . Considérons les points

$$M(k_1 - i + v), M(k_2 - i + v), \dots, M(k_p - i + v), \dots$$

En vertu de la continuité de la fonction ψ , on aura

$$\lim M(k_p - i + v) = M(v) \quad \text{pour } p = \infty.$$

Donc $M(v)$, et, par conséquent, tout point de P appartient à P' .

C. Q. F. D.

La condition $D(P, P') = 0$, traduite dans le langage du Chapitre X, signifie que la trajectoire est instable. En effet, le point $M(o)$ n'appartenant pas à P' , le point mobile ne reviendra jamais dans le voisinage de son point de départ. Il y a exception toutefois lorsque la trajectoire est fermée.

THÉORÈME XX. — *Si l'on a pour toutes les trajectoires $D(P, P') = 0$, il y a certainement un cycle limite, à moins que toutes les trajectoires ne soient fermées.*

Soit, en effet, $N(o)$ un point de P' , dans le voisinage duquel se trouve par définition une infinité de points de P . Soit Q l'ensemble des points $N(i)$, c'est-à-dire des antécédents et des conséquents successifs de $N(o)$.

Il n'est pas possible que, dans tout intervalle compris entre deux

points quelconques de P , il y ait un point de Q ; car, dans le voisinage de $N(o)$, il y a une infinité de pareils intervalles : il y aurait donc une infinité de points de Q , ce qui est impossible en vertu de l'hypothèse

$$D(Q, Q') = 0.$$

Soient donc $M(o)$ et $M(i)$ deux points de P entre lesquels il n'y ait aucun point de Q . Il n'y en aura pas non plus entre $M(i)$ et $M(2i)$; car, si le point $N(k)$ se trouvait entre $M(i)$ et $M(2i)$, le point $N(k-i)$ serait entre $M(o)$ et $M(i)$, puisque la fonction ψ est constamment croissante. Il n'y en aura pas non plus dans l'intervalle compris entre $M(pi)$ et $M(pi+i)$. De plus, si le point $M(i)$ est à droite du point $M(o)$, par exemple, le point $M(pi+i)$ sera à droite de $M(pi)$. Lorsque p croîtra, la seconde coordonnée ω du point $M(pi)$ variera donc toujours dans le même sens; elle ira, par exemple, toujours en croissant, mais elle ne pourra croître indéfiniment, sans quoi $N(o)$ se trouverait dans un des intervalles. Cette coordonnée ω tendra donc vers une certaine limite qui correspondra à un point $P(o)$ du cercle méridien : on aura donc

$$\lim M(pi) = P(o) \quad \text{pour } p = \infty.$$

On en déduit

$$\lim M(pi+k) = P(k),$$

$$\lim M(pi+i) = P(i),$$

$$P(o) = P(i).$$

Ainsi le $i^{\text{ème}}$, conséquent du point $P(o)$, est ce point lui-même. Donc la trajectoire qui passe par ce point est fermée, et il est aisé de voir que c'est un cycle limite.

La condition $D(P, P') = P$ signifie que la trajectoire est stable. En effet, dire que le point de départ $M(o)$ appartient à P' , c'est dire que le point mobile reviendra une infinité de fois dans le voisinage dans ce point.

Si donc on a, pour toutes les trajectoires $D(P, P') = P$, la stabilité est certaine. C'est ce qui arrive, par exemple, dans le second des exemples cités plus haut.

Mais on peut se demander s'il n'arrive pas aussi que l'on ait

D P, P' ; = 0 pour certaines trajectoires et D $(P, P') = P$ pour d'autres ; ce sera là l'origine de toutes les difficultés que nous rencontrerons dans la suite.

Avant d'aller plus loin, je vais établir le lemme suivant :

Soient $M(0)$ et $N(0)$ deux points quelconques, $M(i)$ et $N(i)$ leurs $i^{\text{èmes}}$ conséquents. Si ces deux derniers sont contenus tous deux dans l'intervalle $M(0), N(0)$, je dis qu'il y aura un cycle limite.

En effet, s'il en est ainsi, les deux points $M(2i)$ et $N(2i)$ seront compris tous deux dans l'intervalle $M(i)N(i)$, les deux points $M(3i)$ et $N(3i)$, dans l'intervalle $M(2i)N(2i)$, Donc, lorsque p croîtra, la seconde coordonnée ω du point $M(pi)$ variera toujours dans le même sens, sans pouvoir pourtant dépasser une certaine limite. On en conclut, comme plus haut, l'existence d'un point $P(0)$, tel que

$$\lim M(pi) = P(0) \quad \text{pour } p = \infty,$$

et, par conséquent, celle d'un cycle limite.

Occupons-nous maintenant de déterminer l'ordre circulaire dans lequel sont disposés les points $M(i)$, en laissant de côté le cas où il y a un cycle limite et où cet ordre se détermine aisément.

Soient z_0 la longueur de l'arc $M(0)M(1)$, z_1 celle de l'arc $M(1)M(2)$, ... et, en général, z_i celle de l'arc $M(i)M(i+1)$. Je dis que le rapport

$$\frac{z_i + z_{i+1} + \dots + z_{i+n}}{n}$$

tendra, quand on fera croître n indéfiniment, vers une limite finie, déterminée, indépendante de i , mais incommensurable avec $2\pi r$.

Prenons, à partir du point $M(0)$, un arc L du cercle méridien, égal à k circonférences entières, c'est-à-dire à $2k\pi r$, et supposons

$$(1) \quad z_0 + z_1 + \dots + z_\nu < 2k\pi r < z_0 + z_1 + z_2 + \dots + z_{\nu+1}.$$

D'après cette inégalité, il y a, sur l'arc L , $\nu + 2$ points de P , à savoir : $M(0), M(1), \dots, M(\nu + 1)$, et il n'y en a pas davantage.

Pour bien entendre cette proposition et pour éviter toute confusion, il importe de faire la remarque suivante :

La seconde coordonnée ω d'un point M du cercle méridien $\varphi = 0$ n'est pas entièrement déterminée; car on retrouve le même point en augmentant ω d'un multiple de 2π . Si toutefois on se donne ω_0 , les coordonnées ω_i des conséquents successifs $M(i)$ sera *entièrement* déterminée par les équations

$$\omega_1 = \frac{1}{2} \omega_0, \quad \omega_i = \frac{1}{2} \omega_{i-1}.$$

Nous supposons donc que nous nous sommes donné ω_0 , et nous pourrions regarder ω_i comme complètement déterminé.

Il pourra être avantageux, dans certains cas, d'envisager non pas la coordonnée ω_i elle-même, mais une quantité congrue à ω_i suivant le module 2π et comprise entre z et $z + 2\pi$ (z étant la seconde coordonnée d'un certain point N du cercle méridien). Nous désignerons cette quantité par la notation (ω_i, z) , et nous l'appellerons coordonnée du point $M(i)$ réduite par rapport au point N .

Ainsi, dans la démonstration du théorème XX, nous avons dit que, quand on faisait croître p , la seconde coordonnée ω_{pi} du point $M(pi)$ variait toujours dans le même sens, sans jamais dépasser une certaine limite. Il s'agissait, non de la quantité ω_{pi} elle-même, telle que nous venons de la définir, mais de cette coordonnée réduite par rapport au point $N(0)$.

Au contraire, dans tout ce qui va suivre, il s'agira toujours, sauf avis contraire, de la coordonnée ω_i elle-même.

Ainsi, quand je dis que l'arc L contient les $\nu + 2$ conséquents successifs $M(0), M(1), \dots, M(\nu + 1)$, je veux dire que les coordonnées $\omega_0, \omega_1, \dots, \omega_{\nu+1}$ sont comprises entre ω_0 et $\omega_0 + 2k\pi$.

En écrivant les inégalités (1), j'ai supposé que tous les arcs $z_0, z_1, \dots, z_i$ étaient positifs. C'est, en effet, ce qui arrive. Je puis toujours supposer que

$$\omega_1 > \omega_0,$$

car nous ne pouvons avoir $\omega_1 = \omega_0$, sans quoi nous aurions affaire à une trajectoire fermée, ce que nous ne supposons pas, et, si nous avions $\omega_1 < \omega_0$, nous compterions les arcs ω en sens contraire.

Le deuxième principe donne alors

$$\omega_{i+1} > \omega_i$$

ou

$$z_i > 0,$$

C. Q. F. D.

Soient maintenant $N(o)$ un point quelconque contenu entre $M(o)$ et $M(1)$, ω'_0 sa coordonnée, de telle sorte que

$$\omega_0 < \omega'_0 < \omega_1.$$

Soit ω'_i la coordonnée de son $i^{\text{ème}}$ conséquent $N(i)$. Je dis que

$$\omega_i < \omega'_i < \omega_{i+1};$$

car, si l'on avait

$$\omega'_i < \omega_i,$$

les deux points $N(o)$ et $N(i)$ seraient compris entre les points $M(o)$ et $M(i)$, et si l'on avait

$$\omega'_i > \omega_{i+1},$$

les deux points $M(1)$ et $M(i+1)$ seraient compris entre les points $N(o)$ et $N(i)$. Dans les deux cas, en vertu d'un lemme démontré plus haut, il y aurait un cycle limite, ce que nous ne supposons pas.

De même, si nous avions

$$\omega_k < \omega'_0 < \omega_{k+1},$$

nous en déduirions

$$\omega_{i+k} < \omega'_i < \omega_{i+k+1}.$$

D'ailleurs, nous avons, par l'inégalité (1),

$$\omega_{v+1} < \omega_0 + 2k\pi < \omega_{v+2},$$

et puisque

$$\omega_1 + 2k\pi = \psi(\omega_0 + 2k\pi),$$

$$\omega_{v+2} < \omega_1 + 2k\pi < \omega_{v+3}.$$

Si donc nous prenons, à partir du point $M(1)$, un arc égal à $2k\pi r$, c'est-

à-dire à L , il y aura sur cet arc $\nu + 2$ points de P , à savoir : $M(1)$, $M(2)$, ..., $M(\nu + 2)$.

En raisonnant de même, on verrait que, si l'on prend, à partir d'un point de P , un arc égal à L , il y aura sur cet arc $\nu + 2$ points de P .

Si maintenant on prend un point N sur l'arc α_0 , puis, à partir de ce point N , un arc égal à L , cet arc contiendra certainement les points $M(1)$, $M(2)$, ..., $M(\nu + 1)$, il contiendra ou ne contiendra pas le point $M(\nu + 2)$, et il ne contiendra certainement pas $M(\nu + 3)$.

On raisonnerait de même si le point N était sur un arc quelconque α_i , d'où la conclusion suivante :

Le nombre des points de P situés sur un arc quelconque égal à $2k\pi r$ est égal à $\nu + 1$ ou à $\nu + 2$.

Si l'on donne à k une valeur différente k' , il en résultera pour ν une valeur différente ν' . Je considère les deux intervalles compris, d'une part, entre $\frac{\nu+1}{k}$ et $\frac{\nu+2}{k}$ et, d'autre part, entre $\frac{\nu'+1}{k'}$ et $\frac{\nu'+2}{k'}$. Je dis que ces deux intervalles auront une partie commune. Considérons, en effet, un arc égal à $2kk'\pi r$ et sur lequel il y ait M points de P . Il viendra

$$\begin{aligned} (\nu + 1)k' < M < (\nu + 2)k', \\ \nu' + 1 < k' < M < \nu' + 2 < k', \end{aligned}$$

ce qui démontre la proposition énoncée. Il en serait encore de même si, au lieu de considérer seulement deux valeurs de k et deux intervalles, nous en avions considéré trois ou plusieurs.

Ainsi, si l'on donne à k toutes les valeurs entières positives, tous les intervalles $\frac{\nu+1}{k}$, $\frac{\nu+2}{k}$ auront une partie commune, et, de plus, l'étendue de l'intervalle tendra vers zéro.

Donc $\frac{\nu+1}{k}$ et $\frac{\nu+2}{k}$ tendront vers une limite commune, finie et déterminée que j'appelle μ .

Le nombre μ ne peut pas être commensurable.

En effet, s'il était égal à 2 par exemple, il faudrait que nous eussions

$$\nu + 1 = 2k \quad \text{ou} \quad \nu + 2 = 2k.$$

Soient, pour $k = 1$, $\nu + 1 = 2$, $\nu + 2 = 3$.

On aura, pour $k > 1$,

$$\nu + 2 \geq 2k + 1;$$

d'où

$$\nu + 2 = 2k + 1.$$

On en déduit aisément que l'ordre circulaire des points $M(i)$, où l'indice i est positif, est l'inverse du suivant :

$$[M(0), M(2), M(4), \dots, \text{ad inf.}; M(1), M(3), \dots, \text{ad inf.}]$$

ce qui impliquerait l'existence d'un cycle limite.

Si l'on avait, au contraire, pour $k = 1$, $\nu + 2 = 2$, il viendrait, pour $k > 1$,

$$\nu + 2 \leq 2k;$$

d'où

$$\nu + 2 = 2k.$$

L'ordre circulaire serait alors l'ordre inverse du précédent, et il y aurait encore un cycle limite.

Donc μ est incommensurable.

On aura évidemment

$$\lim_{n \rightarrow \infty} \frac{z_1 + z_{1+n} + \dots + z_{1+n\mu}}{n} = \frac{2\pi i}{\mu} \quad \text{pour } n \rightarrow \infty,$$

ce que nous avons annoncé.

Nous pouvons donc dire que l'ordre circulaire des points de P est caractérisé par un certain nombre incommensurable μ , et c'est ce résultat que nous allons chercher maintenant à énoncer d'une façon plus précise.

Pour mettre en évidence ce fait que ν est fonction de k , j'écrirai $\nu(k)$, au lieu de ν . J'envisagerai ensuite les $\nu(k) + 2$ points :

$$(2) \quad M(0), M(1), M(2), \dots, M[\nu(k) + 1],$$

et cherchons dans quel ordre circulaire ils sont placés.

Il faudra d'abord placer $M(0), M(1), \dots, M[\nu(1) + 1]$ sur le cercle méridien dans l'ordre de leurs indices. Ensuite $M[\nu(1) + 2]$ viendra se

placer entre $M(0)$ et $M(1)$, $M[v(1) + 3]$ entre $M(1)$ et $M(2)$, et ainsi de suite jusqu'à $M[v(2) + 1]$ qui viendra se placer entre $M[v(1) + 1]$ et $M(0)$. Le point suivant $M[v(2) + 2]$ sera placé entre $M(0)$ et $M(1)$; il reste à savoir s'il sera avant $M[v(1) + 2]$ ou après ce point. Il sera avant si

$$v(2) + 2 = 2v(1) + 3$$

et après si

$$v(2) + 2 = 2v(1) + 4,$$

ce qui sont les deux seuls cas possibles.

On continuera de la sorte, et l'on ne sera jamais embarrassé pour placer un nouveau point si l'on connaît les k nombres

$$v(1), v(2), \dots, v(k).$$

Ainsi l'ordre circulaire des points (2) ne dépend que de ces k nombres, c'est-à-dire de μ , et il est le même que celui des quantités $\mu i - E(\mu i)$. Mais k peut être pris aussi grand que l'on veut.

Donc l'ordre circulaire des points $M(i)$ ne dépend que du nombre incommensurable μ , et il est le même que celui des quantités $\mu i - E(\mu i)$.

$E(x)$ désignant, suivant la coutume, le plus grand entier contenu dans x]

Il résulte immédiatement de ce théorème qu'entre deux points quelconques de P il y en a une infinité d'autres; donc, entre deux points quelconques de P , il y aura toujours des points de P' .

Soit N un point quelconque de P' . Dans le voisinage de ce point, il y aura, par définition, une infinité de points de P ; entre deux quelconques de ceux-ci, il y aura toujours un point de P' . Donc, dans le voisinage de N , il y a une infinité de points de P' , c'est-à-dire que N appartient à P'' , ensemble dérivé de P' . On peut donc écrire

$$D(P', P'') = P'.$$

D'autre part, un théorème de la théorie générale des ensembles donne

$$D(P', P'') = P'';$$

d'où

$$P' = P''.$$

Ainsi P' se confond avec son dérivé; c'est donc un de ces ensembles que M. Cantor appelle parfaits.

Mais on sait que M. Cantor distingue les ensembles parfaits linéaires en deux classes : ceux qui ne sont condensés dans aucun intervalle et ceux qui sont condensés dans certains intervalles.

Nous pouvons donc faire trois hypothèses dans l'énoncé desquelles nous reprendrons le langage ordinaire.

1° Nous pouvons supposer que tous les points de la circonférence méridienne $\varphi = 0$ appartiennent à P' .

2° Nous pouvons supposer ensuite qu'il y a sur cette circonférence certains arcs dont tous les points appartiennent à P' , sans que, cependant, il en soit de même de tous les points de cette circonférence.

3° Nous pouvons supposer enfin qu'il n'y a sur cette circonférence aucun arc dont tous les points appartiennent à P' .

Je vais commencer par faire voir que la deuxième hypothèse doit être rejetée. Si on l'adoptait, en effet, on pourrait trouver sur la circonférence un arc AB dont tous les points appartiennent à P' ; mais tel que, si on le prolonge au delà de A ou au delà de B , tous les points du prolongement n'appartiennent pas à P' .

Soient A_i et B_i les $i^{\text{èmes}}$ conséquents de A et B ; les $i^{\text{èmes}}$ conséquents des différents points de l'arc AB seront les différents points de l'arc A_iB_i , et ils devront évidemment faire tous partie de P' .

Je dis que les deux arcs AB et A_iB_i ne peuvent avoir aucune partie commune; car, si A et B étaient situés sur l'arc A_iB_i ou si A_i et B_i étaient situés sur l'arc AB , il y aurait un cycle limite, ce que nous ne supposons pas. Si maintenant le point A_i était situé sur l'arc AB et B_i en dehors de cet arc (ou si B_i était situé sur l'arc AB et A_i en dehors de cet arc), tous les points de l'arc BB_i (ou tous ceux de l'arc AA_i) qui forme un prolongement de l'arc AB au delà du point B (ou du point A) appartiendraient à P' , ce qui est contraire à l'hypothèse faite plus haut.

Soit $M(0)$ un point de P situé sur l'arc AB , le point $M(i)$ sera sur l'arc A_iB_i et par conséquent en dehors de AB . Mais, comme le nombre i est quelconque, il n'y aurait sur l'arc AB aucun conséquent ou antécédent de $M(0)$, c'est-à-dire aucun point de P à l'exception de $M(0)$. Mais, pour que tous les points de l'arc AB appartiennent à P' , il faut qu'il y ait sur cet arc une infinité de points de P .

La seconde hypothèse doit donc être rejetée et il nous reste à examiner la première et la troisième.

La première hypothèse peut certainement être réalisée; car nous avons reconnu qu'elle l'est en effet par les équations

$$\frac{d\omega}{dt} = a, \quad \frac{dz}{dt} = b,$$

lorsque le rapport $\frac{a}{b}$ est incommensurable.

Je dis d'abord que, si tous les points de la circonférence appartiennent à P' , cela sera vrai, quelle que soit la trajectoire à l'aide de laquelle on ait formé les ensembles P et P' . Considérons en effet une autre trajectoire coupant le cercle méridien en un certain nombre de points formant un ensemble Q .

Je dis que le dérivé Q' de Q se composera comme P' de tous les points de la circonférence. En effet soit $N(o)$ un point quelconque de cette circonférence, $N(i)$ son $i^{\text{ième}}$ conséquent, Q l'ensemble des points $N(i)$. Dans le voisinage de $N(o)$ il y a par hypothèse une infinité de points de P , $M(k_1)$, $M(k_2)$, ..., $M(k_n)$.

Lorsque n croît indéfiniment, l'expression $\mu k_n - E(\mu k_n)$ tend vers une certaine limite h . Cette limite caractérise la position du point $N(o)$ sur la circonférence.

Je veux dire que les trois points $M(j)$, $M(k)$ et $N(o)$ se présenteront dans le même ordre circulaire que les trois nombres, $\mu j - E(\mu j)$, $\mu k - E(\mu k)$ et h , quels que soient les indices j et k .

On verrait aisément que la position de $N(i)$ est caractérisée de la même façon par le nombre $k + \mu i - E(h + \mu i)$, et l'on en conclurait que sur tout arc de la circonférence il y a une infinité de points de Q , ce qui revient à dire que tous les points de la circonférence appartiennent à Q' .

Un point quelconque de la circonférence est caractérisé comme on vient de le voir, par un certain nombre h . Ce nombre h se réduit à $\mu i - E(\mu i)$, lorsqu'il s'agit d'un point $M(i)$ appartenant à P et ayant pour coordonnée ω_i . Il se réduit à zéro pour le point $M(o)$ dont la coordonnée est ω_0 . De plus, lorsque ω_i varie de ω_0 à $\omega_0 + 2\pi$, h va constamment en croissant depuis zéro jusqu'à 1. Cela résulte de ce que

les points de la circonférence se succèdent dans le même ordre circulaire que leurs nombres caractéristiques.

Donc h est une fonction croissante de ω , entre les limites ω_0 et $\omega_0 + 2\pi$; on peut donc écrire, en série trigonométrique convergente,

$$2\pi h = \sum A_m \cos m\omega + \sum B_m \sin m\omega$$

ou bien

$$(2) \quad 2\pi h - \omega = \sum A'_m \cos m\omega + \sum B'_m \sin m\omega = \mathcal{G}(\omega).$$

La fonction $\mathcal{G}(\omega)$ représentée par la série (2) sera une fonction continue et périodique de ω . La loi de conséquence pourra alors s'écrire

$$\omega_1 + \mathcal{G}(\omega_1) = \omega_0 + \mathcal{G}(\omega_0) + 2\pi\mu.$$

Soit un point $N(o)$ du cercle méridien ayant pour nombre caractéristique h ; construisons la trajectoire qui passe par $N(o)$ et prolongeons-la jusqu'à ce qu'elle rencontre de nouveau en $N(1)$ le cercle méridien. Si nous attachons à chacun des points de cet arc de trajectoire $N(o)$, $N(1)$ [à l'exception du point $N(1)$], le nombre caractéristique h , tous les points du tore auront un nombre caractéristique et un seul. L'expression

$$2\pi h - \omega + \mu\varphi$$

sera une fonction continue et périodique de ω et de φ , exprimable par une série trigonométrique de la forme suivante :

$$(3) \quad H(\omega, \varphi) = \sum A_{mn} \cos(m\omega + n\varphi + \lambda_{mn}).$$

L'intégrale générale des équations proposées s'écrira alors

$$\omega = \mu\varphi + H(\omega, \varphi) + \text{const.}$$

On est donc certain d'avance que cette intégrale générale peut s'exprimer à l'aide d'une série trigonométrique, mais nous n'avons aucune méthode pour calculer les coefficients de cette série.

Dans le cas qui nous occupe, on a évidemment

$$D(P, P') = P$$

pour toutes les trajectoires et l'on ne peut tracer sur le tore de région si petite qu'elle ne soit traversée une infinité de fois par toutes les trajectoires.

Venons maintenant à la troisième hypothèse, ce qui revient à supposer, comme on le verra, que $D(P, P') = P$ pour certaines trajectoires, et que $D(P, P') = \emptyset$ pour d'autres.

Imaginons donc que P' forme un ensemble parfait qui n'est condensé dans aucun intervalle ou, pour parler le langage ordinaire, que l'on ne puisse trouver sur la circonférence aucun arc dont tous les points appartiennent à P' . Comme P' est un ensemble parfait, on pourra certainement trouver sur la circonférence un arc $A(o)B(o)$, dont les extrémités appartiennent à P' , et dont aucun autre point n'appartient à P' (cf. CANTOR, *Acta mathematica*, t. IV, fasc. 4).

Il arrivera alors que tous les arcs $A(i)B(i)$ jouiront de la même propriété; d'où il résultera que deux arcs $A(i)B(i)$ et $A(k)B(k)$ n'auront aucune partie commune.

Cantor a démontré que, si l'on a sur une circonférence, par exemple, un ensemble parfait de points qui n'est condensé dans aucun intervalle (*loc. cit.*), les points de la circonférence se partageront en trois catégories :

- 1° Les points de certains arcs dont aucun point n'appartient à l'ensemble;
- 2° Les extrémités de ces arcs qui appartiennent évidemment à l'ensemble;
- 3° Enfin les points de l'ensemble qui ne sont pas les extrémités de pareils arcs.

Dans le cas qui nous occupe, les arcs dont les points sont en dehors de l'ensemble sont les arcs $A(i)B(i)$, les extrémités de ces arcs seront les points $A(i)$ et $B(i)$ eux-mêmes; enfin les autres points de l'ensemble P' s'obtiendront en cherchant les points dans le voisinage desquels il y a une infinité de points $A(i)$ et $B(i)$.

Il y aura donc aussi trois espèces de trajectoires :

1° Je citerai en premier lieu les deux trajectoires qui passent, l'une par le point $A(o)$, l'autre par $B(o)$ et que j'appellerai les deux trajectoires A et B . Pour ces deux trajectoires, on a évidemment

$$D(P, P') = P.$$

2° Viennent ensuite les trajectoires qui rencontrent un des arcs $A(i)B(i)$ et qui par conséquent les rencontrent tous. Si le point $C(o)$ d'une de ces trajectoires est sur l'arc $A(o)B(o)$, son $i^{\text{ème}}$ conséquent $C(i)$ sera sur l'arc $A(i)B(i)$. Dans le voisinage d'un point quelconque de P' , il y a une infinité de points $A(i)$ et $B(i)$, il y aura donc aussi une infinité de points $C(i)$. Si donc Q désigne l'ensemble des points $C(i)$ et Q' son dérivé, on aura

$$Q' = P'.$$

Il est clair d'ailleurs que

$$D(Q, Q') = o.$$

3° Il y a enfin des trajectoires qui passent par les points de la troisième catégorie et qui, par conséquent, ne rencontrent aucun des arcs $A(i)B(i)$. Pour ces trajectoires, l'ensemble P' est encore le même et l'on a d'ailleurs

$$D(P, P') = P.$$

Ainsi l'ensemble P' est le même pour toutes les trajectoires, et l'on a

$$D(P, P') = P \quad \text{ou} \quad o,$$

suivant la trajectoire considérée.

Les arcs $A(i)B(i)$ se succèdent dans le même ordre circulaire que les nombres $\mu i - E(\mu i)$.

Les formules

$$\omega_1 + \theta(\omega_1) = \omega_0 + \theta(\omega_0) + 2\pi\mu,$$

$$\omega = \mu\varphi + H(\omega, \varphi) + \text{const.},$$

où $\theta(\omega)$ et $H(\omega, \varphi)$ désignent des séries trigonométriques représentant des fonctions continues et périodiques, subsistent encore ici, mais elles n'ont plus la même portée. En effet, la fonction $\omega + \theta(\omega)$ reste constante tout le long de l'arc $A(i)B(i)$, et l'on trouverait de même sur le tore des régions où la fonction $H(\omega, \varphi) + \mu\varphi - \omega$ conserve une valeur constante.

Il resterait à voir si cette troisième hypothèse, dont nous venons de développer quelques conséquences, peut se réaliser, ou, en d'autres termes, si elle est compatible avec les trois principes que nous avons énoncés plus haut au sujet de la loi de conséquence,

$$\omega_1 = \psi(\omega_0)$$

et avec la forme particulière des équations différentielles considérées.

Je puis affirmer qu'elle est compatible avec les deux premiers principes, en vertu desquels la fonction ψ est continue et croissante.

Est-elle également compatible avec le troisième principe, en vertu duquel la fonction ψ est holomorphe? C'est ce qui resterait à examiner.

Il faudrait, ou bien trouver un exemple où la troisième hypothèse soit réalisée, ce que je n'ai pu faire jusqu'ici, ou bien en démontrer l'impossibilité dans tous les cas.

Je n'ai pu non plus arriver à ce résultat, et je crois, d'ailleurs, que l'hypothèse est effectivement réalisable, mais il y a des cas particuliers où l'on peut démontrer qu'il n'en est pas ainsi et qu'on doit s'en tenir à la première hypothèse.

Soient $C(0)$ un point de la circonférence méridienne, $C(1)$, $C(2)$, ..., $C(i)$ ses $i^{\text{èmes}}$ premiers conséquents; on s'arrêtera lorsqu'on arrivera à un $(i+1)^{\text{ème}}$ conséquent situé sur l'arc $A(0)A(1)$. Soient maintenant M_0 et m_0 , M_1 et m_1 , ..., M_{i-1} et m_{i-1} , M_i et m_i la plus grande et la plus petite valeur que peut prendre la dérivée $\frac{d\omega_1}{d\omega_0}$, respectivement sur l'arc $C(0)C(1)$, sur l'arc $C(1)C(2)$, ..., sur l'arc $C(i-1)C(i)$ et enfin sur l'arc $C(i)C(0)$. Si

$$M_0 M_1 M_2 \dots M_{i-1} M_i < 1 \quad \text{et} \quad M_0 M_1 M_2 \dots M_{i-1} < 1$$

ou bien encore si

$$m_0 m_1 m_2 \dots m_{i-1} m_i > 1 \quad \text{et} \quad m_0 m_1 m_2 \dots m_{i-1} > 1,$$

on est certain que la troisième hypothèse doit être rejetée.

Par le même procédé on peut trouver, dans tous les cas, la limite supérieure de l'un quelconque des arcs $A(i)B(i)$, définis plus haut. Si M est le maximum et m le minimum de la dérivée $\frac{d\omega_1}{d\omega_0}$, tous les arcs $A(i)B(i)$ sont plus petits que

$$\frac{2\pi}{\frac{1}{1-M} + \frac{1}{1-m}}.$$

On peut faire encore une remarque; reprenons le point $A(o)$ et la trajectoire A qui passe par ce point. Soient $C(o)$ et $D(o)$ deux points très voisins de $A(o)$, mais situés l'un à droite et l'autre à gauche de ce point; soient C et D les trajectoires qui passent par ces points. Nous pouvons imaginer trois points mobiles a , c , d décrivant ces trois trajectoires simultanément, de façon à se trouver toujours tous les trois sur le même cercle méridien. Lorsque le temps croîtra, il arrivera alors que la distance ac tendra vers zéro, et que la distance ad ne tendra pas vers zéro, et cela, quelque voisins que soient de $A(o)$ les points $C(o)$ et $D(o)$. Dans la première hypothèse, au contraire, il est impossible que la distance de deux points mobiles a et c décrivant deux trajectoires différentes tende vers zéro. Je ne doute pas qu'on ne puisse se servir utilement de cette remarque, bien que je n'en aie pu moi-même tirer aucun parti.

Toutes ces considérations n'ont pu encore me conduire à la solution de la question principale et ne m'ont pas permis de démontrer rigoureusement que la troisième hypothèse peut effectivement être réalisée.

Bien d'autres questions se posent d'ailleurs, qui sont intimement liées avec la précédente. C'est ainsi qu'on peut se demander comment varie le nombre caractéristique μ , défini plus haut, quand on fait varier les coefficients des équations différentielles. On peut démontrer que cette variation est continue. Mais nous pouvons supposer que, pour certaines valeurs de ces coefficients, il y ait un cycle limite; le nombre μ peut

alors être regardé comme commensurable. Si, pour d'autres valeurs des coefficients, le nombre μ est incommensurable (ou commensurable sans qu'il y ait de cycle limite) et si l'on fait varier ces coefficients d'une manière continue et de façon à passer par les valeurs qui donnent un cycle limite, le nombre μ présentera toujours, pour ces valeurs, soit un maximum, soit un minimum. Il y a là certainement le point de départ d'une série d'études qui seront sans doute fécondes, et que je crois devoir signaler aux travailleurs.

J'espère, d'ailleurs, pouvoir, dans un Chapitre ultérieur de ce Mémoire, donner sur ces questions des résultats plus complets. J'y reviendrai, en effet, après avoir posé, dans la suite de ce travail, une série de problèmes, fort analogues au précédent, mais plus délicats encore, et qui sont liés intimement avec l'importante question de la convergence des séries trigonométriques et, en particulier, des séries employées en Mécanique céleste.

Avant de terminer, je voudrais montrer, en quelques mots, en quoi consiste le lien que je viens de signaler entre le problème qui vient de nous occuper et les méthodes de la Mécanique céleste.

Nous avons vu que l'intégrale générale des équations proposées pouvait se mettre sous la forme

$$\omega - \mu\varphi - H(\omega, \varphi) = \text{const.},$$

H étant une série trigonométrique. Supposons que notre équation différentielle s'écrive

$$\frac{d\omega}{d\varphi} = \mu_0 + \alpha F,$$

μ_0 étant une constante, α un coefficient très petit et F un polynôme en $\cos\omega$, $\sin\omega$, $\cos\varphi$, $\sin\varphi$, ou plutôt encore une série trigonométrique

$$F = \sum A_{mn} \cos(m\omega + n\varphi + \lambda_{mn}).$$

Il viendra

$$(\mu_0 + \alpha F) \left(1 - \frac{dH}{d\omega} \right) = \mu + \frac{dH}{d\varphi}.$$

Supposons que μ et H puissent se développer suivant les puissances de

x , de telle sorte que

$$\begin{aligned} \mu &= \mu_0 + \mu_1 x + \mu_2 x^2 + \dots, \\ H &= H_1 x + H_2 x^2 + H_3 x^3 + \dots \end{aligned}$$

Il viendra

$$\begin{aligned} \mu_0 \frac{dH_1}{d\omega} + \frac{dH_1}{dz} &= F - \mu_1, \\ \mu_0 \frac{dH_2}{d\omega} + \frac{dH_2}{dz} &= \frac{dH_1}{d\omega} F - \mu_2, \\ \mu_0 \frac{dH_3}{d\omega} + \frac{dH_3}{dz} &= \frac{dH_2}{d\omega} F - \mu_3, \\ &\dots\dots\dots \end{aligned}$$

On choisira les constantes $\mu_1, \mu_2, \mu_3, \dots$ de façon que les seconds membres des égalités précédentes soient des séries trigonométriques débarrassées de leur terme tout connu. Il sera alors possible (si μ_0 est incommensurable) de trouver des séries trigonométriques $H_1, H_2, \dots$ satisfaisant *formellement* aux équations précédentes.

Il est impossible de n'être pas frappé de l'analogie de ce procédé d'approximation avec la méthode de M. Lindstedt en Mécanique céleste, et de ne pas comprendre que la question de la convergence du procédé que je viens d'exposer est intimement liée à celle de la convergence des séries employées par le savant astronome de Dorpat. Mais le problème que nous traitons ici est évidemment beaucoup plus simple que les questions analogues de la Mécanique céleste, et, si les difficultés sont de même nature, elles sont moins nombreuses et sans doute plus aisées à surmonter. C'est cette considération qui m'a engagé à insister sur la question qui a fait l'objet de ce Chapitre, et qui m'y fera sans doute revenir à mesure que je trouverai des résultats nouveaux.

Dans la quatrième Partie de ce travail, j'aborderai l'étude des équations différentielles du second ordre. J'abandonne donc momentanément les équations du premier ordre, en me réservant de revenir dans la suite sur les problèmes relatifs à ces équations et restés encore sans solution, en les rapprochant des problèmes analogues qui se poseront au sujet des équations d'ordre supérieur.

Paris, 15 janvier 1885.