

JOURNAL
DE
MATHÉMATIQUES

PURES ET APPLIQUÉES

FONDÉ EN 1836 ET PUBLIÉ JUSQU'EN 1874

PAR JOSEPH LIOUVILLE

ÉDOUARD SELLING

Des formes quadratiques binaires et ternaires

Journal de mathématiques pures et appliquées 3^e série, tome 3 (1877), p. 153-206.

http://www.numdam.org/item?id=JMPA_1877_3_3_153_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>

Des formes quadratiques binaires et ternaires [*];

PAR M. ÉDOUARD SELLING.

IV. — FORMES TERNAIRES INDÉFINIES.

*a. — Relations avec les formes positives correspondantes.**Formes réduites.*

De même que la théorie des formes positives f comprend celle des négatives $-f$, de même aussi la théorie d'une espèce des formes indéfinies f , de laquelle seule je veux traiter, comprend celle des autres $-f$. Pour cela, je suppose que l'invariant des formes f ne soit pas négatif, et provisoirement qu'il ne soit pas zéro. Ces formes f ou $\begin{pmatrix} a, b, c \\ g, h, k \end{pmatrix}$ peuvent ensuite se changer en la forme $x^2 - y^2 - z^2$ par une substitution homogène linéaire à coefficients réels quelconques. Pour ramener la transformation et la réduction de ces formes à celles des formes positives, j'établis le système des équations

$$(9) \begin{cases} \xi^2 - \xi_1^2 - \xi_2^2 = a, & \eta^2 - \eta_1^2 - \eta_2^2 = b, & \zeta^2 - \zeta_1^2 - \zeta_2^2 = c, \\ \eta\zeta - \eta_1\zeta_1 - \eta_2\zeta_2 = g, & \zeta\xi - \zeta_1\xi_1 - \zeta_2\xi_2 = h, & \xi\eta - \xi_1\eta_1 - \xi_2\eta_2 = k, \end{cases}$$

dans lesquelles $\xi, \eta, \zeta, \xi_1, \eta_1, \zeta_1, \xi_2, \eta_2, \zeta_2$ doivent être des variables continues qui parcourent toutes les valeurs réelles possibles. Si l'on détermine par celles-ci, au moyen des équations (7), six quantités variables a, b, c, g, h, k , et si l'on considère une forme $\begin{pmatrix} a, b, c \\ g, h, k \end{pmatrix}$ ou f , qui a ces quantités pour coefficients, cette forme f ou

$$(\xi x + \eta y + \zeta z)^2 + (\xi_1 x + \eta_1 y + \zeta_1 z)^2 + (\xi_2 x + \eta_2 y + \zeta_2 z)^2$$

devient positive pour toutes les valeurs admissibles de ses coefficients variables; je l'appelle *la forme positive correspondant à la forme f* .

[*] Voir, pour la première partie de ce Mémoire, t. III, p. 21 de ce Recueil.

Journ. de Math. (3^e série), tome III. — Mai 1877.

Les équations (9) entraînent les équations

$$(10) \quad \begin{cases} \Xi^2 - \Xi_1^2 - \Xi_2^2 = A, & H^2 - H_1^2 - H_2^2 = B, & Z^2 - Z_1^2 - Z_2^2 = C, \\ HZ - H_1Z_1 - H_2Z_2 = G, & Z\Xi - Z_1\Xi_1 - Z_2\Xi_2 = H, \\ \Xi H - \Xi_1H_1 - \Xi_2H_2 = K, \end{cases}$$

dans lesquelles A, B, C, G, H, K désignent les coefficients de la forme F adjointe à f; car les équations (9), (10) proviennent des équations (7), (8), que je résume par

$$(7a) \quad \begin{vmatrix} a & k & h \\ k & b & g \\ h & g & c \end{vmatrix} = \begin{vmatrix} \xi & \xi_1 & \xi_2 \\ \eta & \eta_1 & \eta_2 \\ \zeta & \zeta_1 & \zeta_2 \end{vmatrix} \cdot \begin{vmatrix} \xi & \eta & \zeta \\ \xi_1 & \eta_1 & \zeta_1 \\ \xi_2 & \eta_2 & \zeta_2 \end{vmatrix},$$

$$(8a) \quad \begin{vmatrix} \Xi & \Xi_1 & \Xi_2 \\ \Xi_1 & \Xi & \Xi_2 \\ \Xi_2 & \Xi_1 & \Xi \end{vmatrix} = \begin{vmatrix} \Xi & \Xi_1 & \Xi_2 \\ H & H_1 & H_2 \\ Z & Z_1 & Z_2 \end{vmatrix} \cdot \begin{vmatrix} \Xi & H & Z \\ \Xi_1 & H_1 & Z_1 \\ \Xi_2 & H_2 & Z_2 \end{vmatrix},$$

si l'on multiplie par le facteur $\sqrt{-1}$ les quantités $\xi_1, \eta_1, \zeta_1, \xi_2, \eta_2, \zeta_2, \Xi_1, H_1, Z_1, \Xi_2, H_2, Z_2$, chaque fois dans les deux systèmes placés à droite de ces équations, ou par le facteur -1 chaque fois dans un des deux systèmes. Et l'on reconnaît l'exactitude des équations (10) à ce que, dans cette transformation, abstraction faite de la modification, ici indifférente, des signes de Ξ, H, Z , les deux systèmes, à droite de l'équation (8a), restent les adjoints des correspondants de (7a), d'où résulte que le système se produisant à gauche de (8a) reste l'adjoint du système se produisant à gauche de (7a) et exprimé par $\begin{vmatrix} a & k & h \\ k & b & g \\ h & g & c \end{vmatrix}$. En même temps, on reconnaît que l'invariant des

formes f et f est le même; je le désigne par I . Si, dans l'une des relations obtenues par (8a), on ajoute des deux côtés, de la manière déjà

employée, le système $\begin{vmatrix} \xi & \xi_1 & \xi_2 \\ \eta & \eta_1 & \eta_2 \\ \zeta & \zeta_1 & \zeta_2 \end{vmatrix}$, on obtient

$$\begin{aligned} \begin{vmatrix} A & K & H \\ K & B & G \\ H & G & C \end{vmatrix} \cdot \begin{vmatrix} \xi & \xi_1 & \xi_2 \\ \eta & \eta_1 & \eta_2 \\ \zeta & \zeta_1 & \zeta_2 \end{vmatrix} &= \begin{vmatrix} \Xi & -\Xi_1 & -\Xi_2 \\ H & -H_1 & -H_2 \\ Z & -Z_1 & -Z_2 \end{vmatrix} \cdot \begin{vmatrix} \Xi & H & Z \\ \Xi_1 & H_1 & Z_1 \\ \Xi_2 & H_2 & Z_2 \end{vmatrix} \cdot \begin{vmatrix} \xi & \xi_1 & \xi_2 \\ \eta & \eta_1 & \eta_2 \\ \zeta & \zeta_1 & \zeta_2 \end{vmatrix} \\ &= \begin{vmatrix} \Xi & -\Xi_1 & -\Xi_2 \\ H & -H_1 & -H_2 \\ Z & -Z_1 & -Z_2 \end{vmatrix} \cdot \begin{vmatrix} \sqrt{I} & 0 & 0 \\ 0 & \sqrt{I} & 0 \\ 0 & 0 & \sqrt{I} \end{vmatrix}. \end{aligned}$$

Des neuf équations comprises là-dedans, trois nous apprennent que $\sqrt{I} \cdot \Xi$, $\sqrt{I} \cdot H$, $\sqrt{I} \cdot Z$ sont les demi-dérivées de $F(\xi, \eta, \zeta)$ relatives à ξ, η, ζ . L'équation $\xi\Xi + \eta H + \zeta Z = \sqrt{I}$ se change donc en

$$(11) \quad F(\xi, \eta, \zeta) = I.$$

Comme pareillement $\sqrt{I} \cdot \xi$, $\sqrt{I} \cdot \eta$, $\sqrt{I} \cdot \zeta$ sont démontrées être les demi-dérivées de $f(\Xi, H, Z)$ relatives à Ξ, H, Z , on a aussi

$$(12) \quad f(\Xi, H, Z) = I.$$

On tire immédiatement les relations entre les coefficients de $f, \mathfrak{f}, F, \mathfrak{F}$ de (11) ou (12) et des équations

$$(13) \quad \begin{cases} a = 2\xi^2 - a, & b = 2\eta^2 - b, & c = 2\zeta^2 - c, \\ g = 2\eta\zeta - g, & h = 2\zeta\xi - h, & k = 2\xi\eta - k, \end{cases}$$

ou

$$(14) \quad \begin{cases} \mathfrak{A} = 2\Xi^2 - A, & \mathfrak{B} = 2H^2 - B, & \mathfrak{C} = 2Z^2 - C, \\ \mathfrak{G} = 2HZ - G, & \mathfrak{H} = 2Z\Xi - H, & \mathfrak{K} = 2\Xi H - K. \end{cases}$$

En effet, les invariants simultanés suivants deviennent

$$(15) \quad Aa + Bb + Cc + 2Gg + 2Hh + 2Kk = -I,$$

$$(16) \quad a\mathfrak{A} + b\mathfrak{B} + c\mathfrak{C} + 2g\mathfrak{G} + 2h\mathfrak{H} + 2k\mathfrak{K} = -I,$$

ce qui résulte aussi sur-le-champ de ce que les invariants des formes $f + f$ et $f - f$ s'évanouissent. Du premier, du déterminant $8\Sigma \pm \xi^2\eta^2\zeta^2$, disparaissent aussi tous les premiers déterminants mineurs. Donc on peut de (15) éliminer encore trois des nombres a, b, c, g, h, k . On obtient, par exemple, au moyen de

$$\begin{aligned} (a + a)(b + b) &= (k + k)^2, & (a + a)(g + g) &= (k + k)(h + h), \\ (a + a)(c + c) &= (h + h)^2, \end{aligned}$$

tant que a n'est pas zéro ou infini,

$$(17) \quad F(a, k, h) = aI,$$

relation identique avec (15).

On obtient de même

$$(18) \quad f(\mathfrak{A}, \mathfrak{B}, \mathfrak{C}) = \mathfrak{A} I.$$

Si f et $\mathfrak{f}$ se changent en f' et $\mathfrak{f}'$ par la substitution $\begin{vmatrix} \alpha & \beta & \gamma \\ \alpha_1 & \beta_1 & \gamma_1 \\ \alpha_2 & \beta_2 & \gamma_2 \end{vmatrix}$,

f' est aussi une forme positive correspondant à f' , comme cela résulte déjà des propriétés d'invariants énoncées.

Les quantités intermédiaires $\xi', \dots$ sont celles qui ont déjà été indiquées; elles peuvent encore être exprimées sommairement par l'équation

$$\begin{vmatrix} \xi' & \xi'_1 & \xi'_2 \\ \eta' & \eta'_1 & \eta'_2 \\ \zeta' & \zeta'_1 & \zeta'_2 \end{vmatrix} = \begin{vmatrix} \alpha & \alpha_1 & \alpha_2 \\ \beta & \beta_1 & \beta_2 \\ \gamma & \gamma_1 & \gamma_2 \end{vmatrix} \cdot \begin{vmatrix} \xi & \xi_1 & \xi_2 \\ \eta & \eta_1 & \eta_2 \\ \zeta & \zeta_1 & \zeta_2 \end{vmatrix},$$

ou

$$\begin{vmatrix} \xi' & \eta' & \zeta' \\ \xi'_1 & \eta'_1 & \zeta'_1 \\ \xi'_2 & \eta'_2 & \zeta'_2 \end{vmatrix} = \begin{vmatrix} \xi & \eta & \zeta \\ \xi_1 & \eta_1 & \zeta_1 \\ \xi_2 & \eta_2 & \zeta_2 \end{vmatrix} \cdot \begin{vmatrix} \alpha & \beta & \gamma \\ \alpha_1 & \beta_1 & \gamma_1 \\ \alpha_2 & \beta_2 & \gamma_2 \end{vmatrix}.$$

Le rapprochement des deux expressions donne

$$\begin{vmatrix} a' & k' & h' \\ k' & b' & g' \\ h' & g' & c' \end{vmatrix} = \begin{vmatrix} \alpha & \alpha_1 & \alpha_2 \\ \beta & \beta_1 & \beta_2 \\ \gamma & \gamma_1 & \gamma_2 \end{vmatrix} \cdot \begin{vmatrix} a & k & h \\ k & b & g \\ h & g & c \end{vmatrix} \cdot \begin{vmatrix} \alpha & \beta & \gamma \\ \alpha_1 & \beta_1 & \gamma_1 \\ \alpha_2 & \beta_2 & \gamma_2 \end{vmatrix},$$

et, si l'on procède aux changements précités de $\xi_1, \eta_1, \zeta_1, \xi_2, \eta_2, \zeta_2$, qui entraînent après eux les changements pareils de $\xi'_1, \eta'_1, \zeta'_1, \xi'_2, \eta'_2, \zeta'_2$, on obtient

$$\begin{vmatrix} a' & k' & h' \\ k' & b' & g' \\ h' & g' & c' \end{vmatrix} = \begin{vmatrix} \alpha & \alpha_1 & \alpha_2 \\ \beta & \beta_1 & \beta_2 \\ \gamma & \gamma_1 & \gamma_2 \end{vmatrix} \cdot \begin{vmatrix} a & k & h \\ k & b & g \\ h & g & c \end{vmatrix} \cdot \begin{vmatrix} \alpha & \beta & \gamma \\ \alpha_1 & \beta_1 & \gamma_1 \\ \alpha_2 & \beta_2 & \gamma_2 \end{vmatrix}$$

Le système des extrémités des longueurs $x(a) + y(b) + z(c)$, ayant un même point de départ, est maintenant variable, et cela de telle sorte que chacun de ses points peut se mouvoir sur un hyperboloïde de révolution équilatère, qui lui est propre, et dont le centre est le point initial, et dont l'axe de révolution est l'axe de ξ, η, ζ . Comme nous ne trouvons, dans les équations (7), que des combinai-

sons des quantités $\xi_1, \eta_1, \zeta_1, \xi_2, \eta_2, \zeta_2$, qui ne changent pas quand le système tourne autour de cet axe, on pourrait admettre que l'une de ces grandeurs, ζ_2 par exemple, est constamment $= 0$; et quand, par hypothèse, $a > 0$, regarder ξ_1, ξ_2 comme des variables indépendantes, dont sont fonctions les huit systèmes de valeurs réelles des autres variables satisfaisant aux équations (9). De deux systèmes de valeurs, dont résultent pour a, b, c, g, h, k les mêmes valeurs, je n'en admettrai qu'un, correspondant au

signe supérieur dans l'équation
$$\begin{vmatrix} \xi & \eta & \zeta \\ \xi_1 & \eta_1 & \zeta_1 \\ \xi_2 & \eta_2 & \zeta_2 \end{vmatrix} = \pm \sqrt{I},$$
 de sorte que

chaque point reste assigné à l'une des deux nappes de son hyperboloïde, si l'hyperboloïde en a deux. Des deux systèmes distincts de valeurs ξ, η, ζ , satisfaisant à l'équation (11), un seulement est admis. Et si l'on veut introduire une nouvelle figure géométrique, on peut regarder ξ, η, ζ comme coordonnées d'un point mobile sur une nappe d'un hyperboloïde à deux nappes.

Je ramène maintenant aussi, pour les formes quadratiques ternaires, la réduction des indéfinies à celle des formes positives, et, *sous la réserve de restrictions ultérieures*, j'appelle *réduite une forme indéfinie*, quand, d'après ma définition, *la forme positive correspondante est réduite pour un système admissible quelconque des valeurs des variables introduites*. On obtient, pour chaque classe, une forme réduite quand, à une forme f quelconque de cette classe, on forme la correspondante f au moyen de n'importe quel système admissible des valeurs de ces variables-là et que l'on applique à f une substitution qui puisse changer f en une réduite. On obtient toutes les formes réduites de cette classe, si l'on agit de même à l'égard de tous les systèmes admissibles de valeurs; de sorte que, dans la première figure et pour $\zeta_2 = 0$, le point ξ, ξ_1, ξ_2 prend toutes les positions sur l'une des nappes de l'hyperboloïde $\xi^2 - \xi_1^2 - \xi_2^2 = a$, ou, dans la seconde figure, que je suivrai désormais, le point ξ, η, ζ prend toutes les positions sur l'une des nappes de l'hyperboloïde $F(\xi, \eta, \zeta) = I$.

*b. — Points angulaires des champs des formes réduites.
Axes de symétrie.*

Ces nappes se divisent en champs tels que, pour tous les points d'un champ, la même forme f' est réduite, tandis que, sur les limites, des coefficients variables de cette forme, g' , h' , k' , l' , m' ou n' sont égaux à zéro. Comme il s'agit seulement de trouver de quelle manière ces champs sont reliés entre eux, et comme cette manière ne change pas, si l'on prend pour base, au lieu de f , une forme équivalente quelconque, on peut, dans tous les cas, prendre pour bases d'autres formes équivalentes, le plus simplement celles dont les formes positives correspondantes sont elles-mêmes réduites et que je désigne aussi par f . Toutefois, j'exclus provisoirement le cas où l'un des nombres a , b , c , d , A , B , C ou $D = F(1, 1, 1)$ deviendrait égal à zéro. Pour reconnaître les champs, on n'a besoin que de considérer les lignes de leurs limites, et pour reconnaître les lignes limites on n'a qu'à considérer leurs points extrêmes, c'est-à-dire les points d'intersection de deux de ces lignes; aucun champ, en effet, ne peut être borné par une ligne rétrogradant sur elle-même; car si, par exemple, cette ligne était $k = 0$, à deux points de cette ligne $\frac{\partial F(\xi, \eta, \zeta)}{\partial \xi}$ ou Z disparaîtrait nécessairement. Il faudrait donc que $C < 0$, et l'équation (11) se changerait en l'équation (4), qui pourtant, avec $k = 0$, n'admet pas deux solutions réelles. Les points à considérer sont de deux espèces essentiellement différentes, suivant que les deux coefficients qui disparaissent sont ou non sur une ligne verticale dans le système $\begin{Bmatrix} g, h, k \\ l, m, n \end{Bmatrix}$. Au point $h = m = 0$, par exemple, suivant III, *b*, alinéa 7, trois champs se rencontrent de telle sorte que, par la substitution $\begin{vmatrix} -1 & 1 & 0 & 0 \\ -1 & 0 & 1 & 0 \\ 0 & 0 & 1 & -1 \end{vmatrix}$, les formes positives correspondantes se changent cycliquement les unes en les autres. Comme lignes limites, entre les trois champs, partent du point $h = m = 0$ les lignes $h = 0$, $h = m$ et

$m=0$, de sorte que chacune de ces lignes n'existe que d'un côté du point. Comme à un point semblable une ligne se fend en deux, j'appelle ce point *un point de fissure*. La condition d'après laquelle on peut obtenir $h'=m'=0$, tandis que f' est réduite, s'exprime facilement par la condition d'après laquelle, dans une forme f , d'où résulte f' par la substitution $\begin{vmatrix} 1 & 0 & 0 & -1 \\ 0 & 0 & -1 & 1 \\ -1 & 1 & 0 & 0 \end{vmatrix}$, on peut obtenir $g=h=k$: donc

$$A \frac{(k+h)(k+k)}{k+g} + B \frac{(k+k)(k+g)}{k+h} + C \frac{(k+g)(k+h)}{k+k} + 2G(k+g) + 2H(k+h) + 2K(k+k) = 2I,$$

tandis que k n'est pas négatif et n'est pas plus grand que a , b ou c , ou, ce qui est la même chose, que, si l'on met λ, μ, ν pour $a-k, b-k, c-k$ ou

$$\frac{(g-h)(g-k)}{k+g} + h+k-a-g, \quad \frac{(h-k)(h-g)}{k+h} + k+g-b-h, \\ \frac{(k-g)(k-h)}{k+k} + g+h-c-k,$$

on obtienne $A\lambda + B\mu + C\nu + Dk = -I$, et que λ, μ, ν, k ne soient pas négatives. Une solution effective de cette équation du quatrième degré en k n'est jamais nécessaire, et même la question de l'existence de racines réelles, en deçà de certaines limites, est vidée, autant que c'est nécessaire, sans calcul par les considérations qu'on verra plus bas. Au point $h=k=0$, par contre, aboutissent simultanément six champs, dont les formes réduites proviennent cycliquement les unes des autres, conformément à ce que nous avons observé

plus haut, à l'alinéa 8, au moyen de la substitution $\begin{vmatrix} 1 & 0 & 0 \\ 0 & 0 & 1 \\ 0 & -1 & 1 \end{vmatrix}$.

Les champs sont séparés les uns des autres dans le même ordre cyclique par les lignes $h=0$, $h+k=0$, $k=0$ et leurs prolongations rétrogrades. J'appelle ces points *points de croisement*. L'équa-

tion (17), avec $h = k = 0$, donne

$$(19) \quad a = \sqrt{\frac{a}{A}} I, \quad b = \frac{k^2}{a + \sqrt{\frac{a}{A}} I} - b = \frac{-C - b \sqrt{\frac{a}{A}} I}{a + \sqrt{\frac{a}{A}} I},$$

$$c = \frac{h^2}{a + \sqrt{\frac{a}{A}} I} - c = \frac{-B - c \sqrt{\frac{a}{A}} I}{a + \sqrt{\frac{a}{A}} I},$$

$$g = \frac{hk}{a + \sqrt{\frac{a}{A}} I} - g = \frac{G - g \sqrt{\frac{a}{A}} I}{a + \sqrt{\frac{a}{A}} I};$$

puis, au moyen de l'équation (18), et de $\mathfrak{H} = \mathfrak{A} = 0$, on obtient encore

$$\mathfrak{A} = \sqrt{\frac{A}{a}} I, \quad \mathfrak{B} = \frac{K^2}{A + \sqrt{\frac{A}{a}} I} - B = \frac{-cI - B \sqrt{\frac{A}{a}} I}{A + \sqrt{\frac{A}{a}} I} \dots$$

Pour que ces valeurs soient réelles, il faut seulement que a et A aient les mêmes signes, et si c'est le négatif, il faut encore, pour que ξ, η, ζ deviennent aussi réelles, que aA soit $\leq I$, condition qui se réalise d'elle-même quand A est positive, comme il résulte de

$$I = aA + kK + hH = aA - (bh^2 - 2ghk + ck^2),$$

où la forme quadratique binaire, placée entre parenthèses, est alors négative. Ici non plus une détermination effective de la racine n'est jamais nécessaire, mais seulement quelquefois un essai pour savoir si elle atteint ou non certaines limites. Il est aisé de voir maintenant que la forme $\begin{Bmatrix} a, b, c \\ g, 0, 0 \end{Bmatrix}$ ou bien est réduite ou se

transforme en une réduite par une substitution $\begin{vmatrix} 1 & 0 & 0 \\ 0 & \beta_1 & \gamma_1 \\ 0 & \beta_2 & \gamma_2 \end{vmatrix}$, suivant

que la forme binaire (b, g, c) est elle-même réduite ou se transforme en une réduite par la substitution $\begin{vmatrix} \beta_1 & \gamma_1 \\ \beta_2 & \gamma_2 \end{vmatrix}$. A la place de (b, g, c) se mettent, en cas de réduction, dans les autres formes ternaires qui sont réduites pour les champs aboutissant au point de croisement, les formes binaires résultant de (b, g, c) par la permutation cyclique des trois coefficients extrêmes et la substitution $\begin{vmatrix} -1 & 0 \\ 0 & -1 \end{vmatrix}$. Chaque système de coefficients entiers $\alpha, \alpha_1, \alpha_2, \lambda, \mu, \nu$ qui satisfont aux conditions $1 = \alpha\lambda + \alpha_1\mu + \alpha_2\nu$ et $0 < f(\alpha, \alpha_1, \alpha_2) \cdot F(\lambda, \mu, \nu) \leq I$ détermine un point de croisement et un seulement.

La question de savoir comment les champs sont reliés les uns aux autres peut maintenant se ramener à la recherche de tous les points de croisement. Il est, à vrai dire, possible qu'un champ n'ait aucun point de croisement, mais seulement des points de fissure à ses limites; toutefois, qu'une ou deux portions de deux lignes, comme $h = 0$ et $m = 0$, forment la limite de ce champ, la ligne $h - m = 0$ seule sera à considérer à l'égard des champs avoisinants. Ces champs se comportent alors comme si les branches respectives de la ligne $h - m = 0$ n'eussent pas été interrompues par les lignes $h = 0$ et $m = 0$.

On doit examiner encore spécialement les cas exceptionnels dans lesquels, conformément à ce qui a été dit plus haut, III, *b*, 9, plusieurs des points étudiés coïncident. Soit, dans un point $g = h = k = 0$, et admettons pour la symétrie que les signes de g, h et k ne soient les mêmes dans aucun des six fragments de surface séparés par les lignes $g = 0, h = 0, k = 0$, ce qui, au besoin, peut être amené par l'emploi d'une des

substitutions $\begin{vmatrix} 1 & 0 & 0 \\ 0 & -1 & 0 \\ 0 & 0 & -1 \end{vmatrix}, \begin{vmatrix} -1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & -1 \end{vmatrix}, \begin{vmatrix} -1 & 0 & 0 \\ 0 & -1 & 0 \\ 0 & 0 & 1 \end{vmatrix}$. Qu'on se figure

ensuite une des trois lignes $g = 0, h = 0$ ou $k = 0$ déplacée d'une quantité infiniment petite, de telle sorte qu'il en résulte un triangle infiniment petit, dans l'intérieur duquel g, h et k soient négatifs et f réduite. Les trois sommets de ce triangle sont alors des points de croisement. Quant aux six champs qui les entourent, ceux qui sont adjacents aux côtés du triangle appartiennent à deux de ces points, tandis que le triangle lui-même appartient à la fois aux trois points. Les deux lignes $g + h = 0$ et $g + k = 0$ bornent le fragment de

surface qui est adjacent au côté du triangle $g = 0$, et dans lequel la forme résultant de f par la substitution $\begin{vmatrix} 1 & 0 & 0 & -1 \\ 1 & 0 & -1 & 0 \\ 0 & 1 & 0 & -1 \end{vmatrix}$ est réduite.

Ces deux lignes se rencontrent en un point de fissure et y ont comme continuation commune la ligne $h = k$. Ce champ devient aussi infiniment petit. Attendu qu'il en est de h et k comme de g , on reconnaît que, si on laisse les points coïncider de nouveau, et qu'on ne tienne pas compte des champs dégénérant en un point, neuf champs se trouvent placés autour du point $g = h = k = 0$ et sont séparés les uns des autres, d'une part par les lignes $g = 0$, $h = 0$, $k = 0$, et de l'autre par les lignes $h = k$, $k = g$, $g = h$. Au même point se restreignent aussi les champs des trois formes encore restantes des seize formes essentiellement différentes et réduites en ce point, les-

quelles résultent de f au moyen des substitutions $\begin{vmatrix} -1 & 0 & 0 & 1 \\ 1 & -1 & 0 & 0 \\ 0 & 0 & 1 & -1 \end{vmatrix}$, $\begin{vmatrix} 1 & 0 & 0 & -1 \\ 0 & -1 & 0 & 1 \\ 0 & 1 & -1 & 0 \end{vmatrix}$ et $\begin{vmatrix} -1 & 0 & 1 & 0 \\ 0 & 1 & 0 & -1 \\ 0 & 0 & -1 & 1 \end{vmatrix}$.

Si deux lignes limites coïncident en une seule, comme, par exemple, quand $g = h = 0$, les lignes $g = 0$ et $h = 0$ se confondant en la ligne $\zeta = 0$, alors la condition prescrite à la forme f est aussi remplie par la forme qui coïncide numériquement avec elle et qui en résulte

au moyen de la substitution $\begin{vmatrix} -1 & 0 & 0 \\ 0 & -1 & 0 \\ 0 & 0 & 1 \end{vmatrix}$; dans ce cas, quatre des neuf

champs généralement finis dégénèrent en la ligne $\zeta = 0$, avec laquelle coïncident aussi les lignes $g + h = 0$ et $g - h = 0$. Sur le côté positif de la ligne $k = 0$ restent les champs des deux formes concordant numériquement l'une avec l'autre et provenant de f par les substitu-

tions $\begin{vmatrix} \mp 1 & 0 & 0 \\ 0 & \pm 1 & 0 \\ 0 & 0 & -1 \end{vmatrix}$; sur le côté négatif, se présentent, de part et

d'autre, les champs des formes concordant numériquement et qui

résultent de f par les substitutions $\begin{vmatrix} \pm 1 & 0 & 0 & \mp 1 \\ 0 & \pm 1 & \mp 1 & 0 \\ 0 & 0 & 1 & -1 \end{vmatrix}$. L'autre ligne

limite de chacun de ces deux champs, passant par le point, est

$\pm h - k = 0$ ou $\mp g - k = 0$, suivant que la valeur absolue de η est plus grande ou plus petite que celle de ξ . Les formes, concordant numériquement entre elles, des champs suivants de part et d'autre et s'étendant jusqu'à la ligne $\zeta = 0$, résultent donc de f ,

ou par les substitutions $\begin{vmatrix} \pm 1 & 0 & 0 & \mp 1 \\ 0 & \pm 1 & 0 & \mp 1 \\ 0 & -1 & 1 & 0 \end{vmatrix}$, ou par les substitutions

$\begin{vmatrix} \mp 1 & 0 & 0 & \pm 1 \\ 0 & \mp 1 & 0 & \pm 1 \\ -1 & 0 & 1 & 0 \end{vmatrix}$. Un seulement de ces deux champs aurait été fini,

dans le cas plus général, l'autre eût dégénéré en un point. Se trouveront réduits à ce point les champs des deux formes provenant de f par

les substitutions $\begin{vmatrix} \pm 1 & 0 & 0 & \mp 1 \\ 0 & \mp 1 & \pm 1 & 0 \\ 0 & 0 & -1 & 1 \end{vmatrix}$; quant à ceux des huit formes res-

tantes, ils dégénèrent en la ligne $\zeta = 0$, savoir, quatre dont les formes ont été nommées ici, sur le côté négatif, et quatre sur le côté positif de la ligne $k = 0$. On reconnaît aussi cette dernière particularité à ce que chaque point de la ligne $\zeta = 0$ a les propriétés d'un point de croisement, où les trois lignes qui se coupent coïncident en une seule, où, par conséquent, quatre des six champs dégénèrent en cette ligne et où les formes des deux champs adjacents peuvent être converties l'une en

l'autre par la substitution $\begin{vmatrix} -1 & 0 & 0 \\ 0 & -1 & 0 \\ 0 & 0 & 1 \end{vmatrix}$.

Comme les formes des champs placés immédiatement des deux côtés de la ligne $\zeta = 0$ concordent numériquement entre elles et que, par la forme appartenant à un champ, les formes de tous les champs voisins sont déterminées complètement, cette ligne forme en quelque sorte un axe de symétrie, des deux côtés duquel, même aux distances plus grandes, la division en champs est pareille; il s'ensuit que, pour reconnaître tous les champs, on n'a pas besoin de dépasser cette ligne. De cette espèce sont, dans la *fig. 1* (*voir la Planche*), qui représente la division en champs pour une classe de formes, à nombres entiers et à l'invariant 60, quatre des six lignes limites externes. On n'y a inscrit que les formes dont les champs ont une étendue finie, en les plaçant à leur intérieur. Les lettres mises aux lignes limites désignent celles des quantités variables qui deviennent égales à zéro, le long de celles-ci, et la désignation est toujours choisie par rapport à la forme, dans

le champ de laquelle la lettre est placée. Ce que ξ , η , ζ sont par rapport à a , b , c , c'est ω par rapport à d . Dans chacune de ces formes, pour exclure l'arbitraire, l'ordre des coefficients a , b , c , d est choisi de telle sorte qu'ils forment une progression décroissante et qu'aussi, parmi les coefficients g , h , k , l , m , n , chaque fois le premier, non encore déterminé, devienne aussi grand que possible. Mais, pour les substitutions indiquées dans le texte, à l'égard desquelles on pourrait désirer différents ordres des a , b , c , d correspondant aux différentes lignes limites d'un champ, il est admis que ces ordres ont d'abord été modifiés de part et d'autre, en cas de besoin.

Si, aux conditions dernièrement étudiées, $g = h = 0$, se joint encore $k = 0$, et si η est le facteur, devenant zéro, de $\mathfrak{k}$, de sorte que l'une des deux branches de la ligne $\mathfrak{g} = 0$ coïncide avec $\mathfrak{h} = 0$, l'autre, avec $\mathfrak{k} = 0$, les conditions prescrites à la forme f sont remplies par quatre formes concordant numériquement entre elles, lesquelles proviennent

les unes des autres par les substitutions $\begin{vmatrix} 1 & 0 & 0 \\ 0 & -1 & 0 \\ 0 & 0 & -1 \end{vmatrix}$, $\begin{vmatrix} -1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & -1 \end{vmatrix}$

et $\begin{vmatrix} -1 & 0 & 0 \\ 0 & -1 & 0 \\ 0 & 0 & 1 \end{vmatrix}$, et dont les champs dégénèrent chacun en systèmes de deux lignes, savoir : d'une portion de la ligne $\zeta = 0$, et d'une portion de la ligne $\eta = 0$.

Une étendue finie n'est donnée qu'aux quatre champs compris chacun entre deux de ces portions de lignes et dont les formes, qui d'ailleurs concordent numériquement entre elles, proviennent des

quatre susdites par la substitution $\begin{vmatrix} 1 & 0 & 0 & -1 \\ 0 & 1 & 0 & -1 \\ -1 & 0 & 1 & 0 \end{vmatrix}$. A cette espèce

appartiennent trois des angles externes de la *fig. 1*.

Si, par contre, et sans avoir $k = 0$, dans le cas $g = h = 0$, la ligne $\zeta = 0$ est coupée par la ligne $\eta = 0$, l'autre partie de la ligne $\mathfrak{g} = 0$, les formes de deux champs séparés par la ligne $\zeta = 0$ se chan-

gent, par la substitution $\begin{vmatrix} -1 & 0 & 0 & 1 \\ 0 & -1 & 0 & 1 \\ -1 & 0 & 1 & 0 \end{vmatrix}$, en les formes de deux autres

champs qui avaient jusque-là dégénéré en cette ligne. De tels points sont placés sur la ligne limite à droite et à l'extrémité inférieure de la ligne limite gauche de la *fig. 1*.

A l'aide des substitutions énoncées, III, b , non-seulement des formes positives réduites, mais encore des formes quadratiques ternaires quelconques se transforment en elles-mêmes, quand elles remplissent les conditions indiquées, et toutes ces conditions peuvent être remplies par des formes indéfinies, excepté celles qui, pour x et λ , donnent la valeur 24. Si les conditions énoncées, de 1 à 4, sont satisfaites dans des formes indéfinies réduites, on trouve, aux lignes limites de leurs champs, une symétrie en x sections, et, par suite de cela, aussi dans toute l'ordonnance des champs plus ou moins avoisinants. Au cas 1, $a = b$, $g = h$, $l = m$, appartiennent, dans la *fig. 1*,

les formes $\begin{vmatrix} -1 & 2 & -3 & 2 \\ & -1 & -3 & 2 \\ & & -2 & 8 \\ & & & -12 \end{vmatrix}$ et $\begin{vmatrix} -1 & 2 & -2 & 1 \\ & -1 & -2 & 1 \\ & & -12 & 16 \\ & & & -18 \end{vmatrix}$. L'axe de la

symétrie à deux sections est indiqué par la ligne ponctuée. Au cas 2,

$a = b$, $c = d$, $g = l$, $h = m$, appartient la forme $\begin{vmatrix} -1 & -2 & -1 & 4 \\ & -1 & 4 & -1 \\ & & -9 & 6 \\ & & & -9 \end{vmatrix}$

de la *fig. 1*. Si les conditions énoncées, de 6 à 9, sont remplies dans des formes indéfinies réduites, des lignes limites déterminées peuvent alors aussi se montrer comme axes de la symétrie. Au cas 6 appar-

tient la forme $\begin{vmatrix} 3 & -5 & -1 & 3 \\ & -1 & 3 & 3 \\ & & -2 & 0 \\ & & & -6 \end{vmatrix}$ de la *fig. 1*, dans laquelle $n = 0$,

$m = l$, où donc, par analogie avec le cas $h = 0$, $g = n$, on a

$\lambda = 2x = 2$, et la forme $\begin{vmatrix} 3 & -7 & 1 & 3 \\ & 3 & 1 & 3 \\ & & -2 & 0 \\ & & & -6 \end{vmatrix}$, dans laquelle $n = 0$,

$m = l$, $g = h$, où donc, par analogie avec le cas $h = 0$, $g = n$,

$k = l$, on a $\lambda = 2x = 4$, de plus, les formes $\begin{vmatrix} -1 & -3 & 3 & 1 \\ & -2 & 0 & 5 \\ & & -6 & 3 \\ & & & -9 \end{vmatrix}$

et $\begin{vmatrix} 3 & 1 & -3 & -1 \\ & -2 & 0 & 1 \\ & & -6 & 9 \\ & & & -9 \end{vmatrix}$, pour lesquelles $\lambda = 2x = 2$.

Au cas 7 appartiennent les formes

$$\begin{vmatrix} 3 & -4 & 0 & 1 \\ & 3 & 1 & 0 \\ & & -9 & 8 \\ & & & -9 \end{vmatrix} \text{ et } \begin{vmatrix} 3 & 1 & -4 & 0 \\ & -2 & 0 & 1 \\ & & -4 & 8 \\ & & & -9 \end{vmatrix},$$

pour lesquelles $\lambda = 2$; la première concorde avec elle-même, après modification convenable de l'ordre des a, b, c, d , la deuxième avec une troisième, dont le champ a un point de fissure commun avec ceux des deux premières.

Les cas 8, $h = k = 0$, auxquels appartient, pour $\lambda = 4$, la forme

$$\begin{vmatrix} -1 & 2 & -1 & 0 \\ & -1 & -1 & 0 \\ & & -18 & 20 \\ & & & -20 \end{vmatrix}, \text{ et 9, } g = h = k = 0, \text{ ont déjà été discutés plus}$$

haut.

Nous avons ainsi épuisé les cas dans lesquels deux formes, concordant numériquement l'une avec l'autre, peuvent appartenir au même champ ou à deux champs contigus l'un à l'autre le long d'une ligne.

c. — Manières de trouver les champs.

Pour trouver, à partir d'un point de croisement $h = k = 0$, les autres du même champ, il faut d'abord rechercher s'il y a des points de croisement $h = l = 0$, $h = g = 0$, $h = n = 0$, ce qui exige que a et $cd - n^2$, c et $ab - k^2$, c et $ad - l^2$ aient les mêmes signes, et, quand ils sont négatifs, ils aient un produit n'excédant pas l'invariant; enfin que les formes binaires (c, n, d) , (a, k, b) , (a, l, d) à considérer suivant (19) soient réduites respectivement. Des points de croisement que l'on trouve réellement, il n'y en peut guère avoir qu'un seul avec $h = k = 0$ sur la même partie de la ligne $h = 0$, et je ne veux considérer comme tel que le point $h = g = 0$, attendu qu'il en serait de même des deux autres, dans le cas présent. Car les deux lignes $h = 0$, $g = 0$ ne peuvent pas, d'après (19), avoir plus d'un point d'intersection réel et si la partie à étudier de la ligne $h = 0$ est entrée, auprès d'un tel point, dans la région de la nappé hyperbo-

loïdale, où $g > 0$, il n'en serait ainsi qu'autant qu'elle rentrerait dans la région où l'on a $g < 0$. Tous les points situés sur une branche de la ligne $h = 0$ se distinguent de ceux qui sont situés sur l'autre par les signes de ξ et ζ , dans le cas où une nappe de l'hyperboloïde (11) est coupée par deux nappes de l'hyperboloïde $2\xi\zeta = h$, savoir, quand B est positif et par conséquent a et c négatifs; et ils se distinguent par les signes de H dans le cas où une nappe de l'hyperboloïde (11) est coupée seulement par une nappe de l'hyperboloïde $2\xi\zeta = h$, mais en deux parties, ce qui a lieu quand B , a et c sont négatifs. Quand les signes de a et c sont différents, la ligne $h = 0$ n'a plus qu'une seule branche; si tous deux sont positifs, cette ligne est impossible, comme il suit de (17). Si le point $h = g = 0$ remplit, comparativement avec le point $h = k = 0$, les conditions indiquées, on doit le considérer comme le point de croisement consécutif, et il faut étudier d'abord la ligne $g = 0$, et ensuite la ligne $h = 0$. Et si, entre les deux points de croisement considérés, la ligne $h = 0$ est coupée deux ou quatre fois par la ligne $m = 0$, on peut n'en tenir aucun compte. Car la ligne $m = 0$ ne peut être coupée plus d'une fois par aucune des autres lignes $g = 0$, $k = 0$, $l = 0$ ou $n = 0$. Une branche de cette ligne, qui entre dans le champ de f et le quitte à des points de la ligne $h = 0$, où g , k , l , n sont négatifs, ne peut, par conséquent, entre les deux points d'intersection, être coupée par aucune de ces quatre lignes; seule, par conséquent, ou à l'aide de l'autre branche, elle ne découpera du champ admis de f qu'un fragment qui, sur sa limite, ne contient pas de points de croisement, et dont, par conséquent, on peut ne pas tenir compte. Si, par cette portion de la ligne $h = 0$, entrent deux branches de la ligne $m = 0$, qui coupent aussi d'autres lignes limites, on peut se figurer les points d'intersection des deux lignes rapprochés jusqu'à ce qu'ils se touchent l'un l'autre, puis les portions situées d'un côté de la ligne $h = 0$ réunies ensemble et détachées de la ligne $h = 0$, sans qu'il y ait un changement essentiel apporté au groupement des champs; seulement, les deux champs séparés seront réunis; quant au champ intermédiaire qui les séparait, il sera lui-même partagé en deux.

Mais si, avec le point $h = k = 0$, sur la même branche de la ligne

$h = 0$, il n'y a pas de point de croisement, pour lequel f est réduite, cette ligne se terminera par un point de fissure et aussi il n'importera pas qu'elle soit coupée une ou trois fois par la ligne $m = 0$. La continuation de la limite du champ de f est donc donnée par une portion de la ligne $m = 0$. Si cette ligne ne porte aucun point de croisement, elle ne fait que réunir deux branches de la ligne $h = 0$, que l'on peut alors traiter comme si elles étaient réunies. Mais, si cette branche porte un point de croisement, on peut partir de ce point comme du point $h = k = 0$. Il ne se présente de difficulté que lorsqu'il y a deux branches de la ligne $h = 0$ et deux de la ligne $m = 0$, dont chacune n'a qu'un point de croisement; car alors on peut se demander comment les dernières s'apparient avec les premières. On pourrait obtenir une réponse en considérant les lignes $\xi = 0$ et $\zeta = 0$, ou $H = 0$, déjà indiquées, et passant entre les deux branches de la ligne $h = 0$, en considérant spécialement les points d'intersection de ces lignes, avec la ligne $m = 0$, desquels, s'ils étaient réels, resterait à décider sur quelles branches de la ligne $m = 0$ ils se trouvent. Suivant que, dans un pareil point, f est réduite ou non, il est situé ou non entre le point de croisement et le point de fissure de la branche respective de la ligne $m = 0$. Mais, comme à la forme f il faut toujours appliquer la même substitution, pour obtenir la forme f' du champ séparé du sien, par la ligne $h = 0$, la forme f' a deux champs distincts, quand les limites sont des portions des branches de la ligne $h = 0$. Si l'on se figure les deux points de fissure en question amenés à coïncider, les deux branches de la ligne $h = 0$ réunies entre elles, ainsi que les deux branches de la ligne $m = 0$, puis la première détachée de la seconde, les deux champs séparés de la forme f' se trouvent réunis; mais le champ de la forme f est partagé en deux, tandis que rien n'est changé aux autres lignes limites. On peut donc, toutes les fois qu'il y a précisément deux points de croisement sur $h = 0$, procéder simplement comme s'ils appartenaient à la même branche de cette ligne. Quand on aura obtenu tous les points de croisement du champ ou des champs d'une forme, il faut que de soi-même la dernière ligne limite se rattache à la première.

Au lieu de déterminer toutes les lignes limites par rapport aux

champs, on peut aussi déterminer tous les champs adjacents par rapport aux lignes limites. Évidemment, d'après le procédé susdit, toute ligne limite commence et finit à un point de fissure et passe par un ou plusieurs points de croisement. Une ligne $h = 0$ ne peut aller à l'infini par un nombre illimité de points de croisement que si h et un des nombres g, k, l, n sont égaux à zéro.

La division, en champs, de la nappe de l'hyperboloïde, aurait aussi pu se baser sur la réduction des formes $\mathcal{F}$ au lieu de se baser sur la réduction des formes f . Les points de croisement des deux divisions coïncident; sur d'autres points, les lignes limites des deux systèmes ne peuvent se couper. Des parties de chaque champ d'un système peuvent appartenir aux champs de trois formes de l'autre système. Si l'on a trouvé la division d'un système, celle de l'autre s'offre d'elle-même: chacune des six lignes limites partant d'un point de croisement dans l'un des systèmes est située entre deux des six de l'autre. Quelques lignes limites du deuxième système, que l'on pourrait appeler *système adjoint*, sont marquées (*fig. 1*) comme lignes serpentantes dans une intention qui sera expliquée ultérieurement: c'est pour le même motif que quelques lignes ordinaires ont été mises en relief.

Nous avons vu que celles des formes binaires, par lesquelles on peut rationnellement représenter le nombre zéro et dont l'invariant est conséquemment un carré négatif, si les coefficients sont rationnels, se distinguent des autres en ce que, chez les premières, sur la branche hyperbolique, les intervalles de certaines formes s'étendent à l'infini, mais non chez les dernières. De même les formes ternaires, qui représentent rationnellement zéro, et pour le critérium desquelles je renvoie aux *Disquisitiones arithmeticae*, art. 299, se distinguent des autres en ce que, chez les premières, les champs de certaines formes s'étendent à l'infini, mais non chez les dernières. Comme dans ces dernières aucun des nombres a, b, c, d ne peut devenir infiniment petit, aucun non plus ne doit acquérir une grandeur infinie, les conditions de réduction étant maintenues. Il ne peut y avoir qu'un nombre fini de formes réduites, à nombres entiers, d'un invariant donné; à plus forte raison ne peut-il y avoir qu'un nombre fini d'une classe: c'est ce qui résulte de ce qu'il n'y a qu'un nombre fini de systèmes de valeurs entiers, de a, b, c, g, h, k , pour lesquels la forme f , constituée d'après (19) ou au moyen de $\xi = 0$, est ré-

duite. On obtient la limitation pour a et A par $0 < aA \leq I$, puis, quand b et c sont négatifs, pour b, c, g, h, k par les conditions de réduction pour (b, g, c) ; quand b et c ont des signes différents, pour b, c, g par $A = bc - g^2$, pour h et k par les conditions de réduction pour (b, g, c) . Si b et c sont positifs, B et C sont négatifs, par conséquent B, C, G, H, K sont limités et par là aussi b, c, g, h, k .

Si donc on pousse suffisamment loin le développement décrit des champs, on peut, tandis que leur nombre est encore fini, arriver à ce que tout champ à ajouter nouvellement ne soit que la répétition d'un champ déjà existant, de sorte que tout le développement ultérieur n'est qu'une répétition à l'infini de ce qui a déjà été effectué. De plus on n'est pas forcé de dépasser, lors du premier développement, des axes de symétrie, tels que ceux qui ont été définis plus haut. Tout axe de ce genre est évidemment illimité dans ses deux directions, attendu que de la symétrie de ses deux côtés, existant en un endroit, découle d'elle-même la continuation de cette symétrie. Dans l'exemple achevé (fig. 1, précitée), toutes les limites du système développé des champs sont de pareils axes de symétrie. Dans cet exemple, il y a une quadruple symétrie à leurs points d'intersection, ce que j'ai voulu indiquer en choisissant des angles droits choisis.

Dans ce qu'on vient de dire, on a montré la possibilité de résoudre tous les problèmes posés dans l'introduction, pour toute forme donnée, comme il est aisé de voir qu'une transformation de f en elle-même correspond à chaque champ, dont la forme concorde numériquement avec f , la substitution pouvant être trouvée à l'aide des substitutions consécutives, par lesquelles la forme d'un champ se change en celle d'un champ avoisinant. Il nous reste à parler des simplifications très-considérables qui naissent de la réunion des champs, dans les formes desquels a et A sont identiques, par analogie avec la réunion des intervalles en espaces, dans les formes binaires, et à faire ressortir les propriétés générales des résultats.

d. — Réunion des points de croisement en territoires. Développement immédiat de ces territoires.

Si dans un système de valeurs admissible quelconque ξ, η, ζ peut se présenter le point de croisement $h = k = 0$, si, par conséquent, la ligne droite (a) peut être perpendiculaire au plan (A) , on peut exprimer l'ensemble des lignes droites, qui peuvent devenir perpendiculaires au plan (A) , par $[f(1, \alpha_1, \alpha_2)]$ ou $(a) + \alpha_1(b) + \alpha_2(c)$, quand α_1, α_2 désignent les nombres entiers, au moyen desquels $f(1, \alpha_1, \alpha_2)$ obtient le signe de A et ne donne pas avec A un produit dépassant l'invariant I , ce qui serait possible si A avait une valeur négative. Si l'on fait partir du point initial de la ligne (a) toutes les lignes $(a) + \alpha_1(b) + \alpha_2(c)$, déterminées comme fonctions de ξ, η, ζ aussi quant à leur longueur, si l'on fait, par conséquent, partir les lignes $\alpha_1(b) + \alpha_2(c)$ du point final de la ligne (a) , alors les points finaux de toutes ces lignes se trouveront dans un plan (A) et, si $A > 0$, dans l'intérieur d'une certaine ellipse; si $A < 0$, entre une certaine hyperbole et ses asymptotes, comme le montrent les conditions respectives

$$1 \leq f(1, \alpha_1, \alpha_2) \quad \text{ou} \quad 1 - \frac{I}{A} \leq (b, g, c) \left(\alpha_1 - \frac{K}{A}, \alpha_2 - \frac{H}{A} \right)$$

et

$$-1 \leq f(1, \alpha_1, \alpha_2) \leq \frac{I}{A} \quad \text{ou} \quad \frac{I}{-A} - 1 \leq (b, g, c) \left(\alpha_1 - \frac{K}{A}, \alpha_2 - \frac{H}{A} \right) \leq 0.$$

Je continue d'exclure provisoirement celles des formes par lesquelles zéro peut être représenté. Le système particulier de valeurs des variables ξ, η, ζ , pour lequel une ligne $[f(1, \alpha_1, \alpha_2)]$ est perpendiculaire sur (A) , par conséquent la position du point de croisement respectif sur la nappe de l'hyperboloïde $F(\xi, \eta, \zeta) = I$, est, comme l'indique

l'emploi de la substitution $\begin{vmatrix} 1 & 0 & 0 \\ \alpha_1 & 1 & 0 \\ \alpha_2 & 0 & 1 \end{vmatrix}$ ou de son adjointe, déterminée par les équations

$$h + g\alpha_1 + c\alpha_2 = k + b\alpha_1 + g\alpha_2 = 0 \quad \text{ou} \quad A - A\alpha_1 = H - A\alpha_2 = 0.$$

Pour l'élimination de α , et α_2 des conditions énoncées, on obtient, quand $A > 0$, $\mathfrak{A}^2 \leq f(\mathfrak{A}, \mathfrak{A}, \mathfrak{H})$ ou, au moyen de l'équation (18) $\mathfrak{A}^2 \leq AI$, par conséquent $2\mathfrak{E}^2 \leq A + \sqrt{AI}$; mais, si $A < 0$, $\mathfrak{A}^2 \leq -AI \leq \frac{I}{A} \mathfrak{A}^2$, par conséquent $A + \sqrt{-AI} \geq 2\mathfrak{E}^2 \geq 0$. Tous ces points de croisement sont donc situés, dans le premier cas, sur un fragment fini de la nappe hyperboloïdale, séparé par un plan et borné par une ellipse; dans le second cas, sur un fragment infini de la nappe hyperboloïdale, renfermé entre deux plans et limité par deux branches d'hyperboles. Leur nombre est, dans le premier cas, évidemment fini; mais, dans le deuxième cas, infiniment grand, de telle sorte qu'un nombre fini de valeurs numériques $f(1, \alpha_1, \alpha_2)$ se répète pour une infinité de systèmes de valeurs α_1, α_2 . On obtient, en effet, un nombre infini de répétitions de la valeur α elle-même, dont chacune produit aussi une répétition des autres

valeurs, quand on applique à f d'abord des substitutions $\begin{vmatrix} \pm 1 & 0 & 0 \\ 0 & \beta_1 & \gamma_1 \\ 0 & \beta_2 & \gamma_2 \end{vmatrix}$ au déterminant 1, par lesquelles (b, g, c) se change en $\pm(b, g, c)$,

puis des substitutions $\begin{vmatrix} 1 & 0 & 0 \\ \beta & 1 & 0 \\ \gamma & 0 & 1 \end{vmatrix}$, par lesquelles quand, indépendam-

ment des signes précités $\pm$, on pose le nombre $\varepsilon = \pm 1$, les nombres $\pm \varepsilon h$ et $\pm \varepsilon k$ se mettent à la place de h et k et, par conséquent, εH et εK à la place de H et K . Si s est le plus grand commun diviseur de b, g, c et $s\sigma$ celui de $b, 2g, c$, si nous avons $t^2 + \frac{A}{s^2}u^2 = \pm \sigma^2$ et

si $\beta_1 = \frac{t}{\sigma} - \frac{gu}{s\sigma}$, $\gamma_1 = \frac{-cu}{s\sigma}$, $\beta_2 = \frac{bu}{s\sigma}$, $\gamma_2 = \frac{t}{\sigma} + \frac{gu}{s\sigma}$, les équations $\varepsilon K = -H\gamma_1 + K\gamma_2 - A\beta$, $\varepsilon H = H\beta_1 - K\beta_2 - A\gamma$, fournies par l'ap-

plication à F de la substitution $\begin{vmatrix} \pm 1 & \mp \beta & \mp \gamma \\ 0 & \pm \gamma_1 & \mp \beta_2 \\ 0 & \mp \gamma_1 & \pm \beta_1 \end{vmatrix}$, donnent

$$A\beta = K \left(\frac{t}{\sigma} - \varepsilon \right) - \frac{Ahu}{s\sigma}; \quad A\gamma = H \left(\frac{t}{\sigma} - \varepsilon \right) + \frac{Aku}{s\sigma}$$

Quand $s = 1$ et qu'on emploie le signe supérieur, quand, par conséquent, on a $(t - \sigma)(t + \sigma) \equiv 0 \pmod{A}$, les côtés droits de ces équations-là sont toujours divisibles par A pour l'une des deux valeurs de ε ,

quand A est un nombre premier ou une puissance d'un nombre premier impair; en outre, il en est de même, quel que soit le signe employé, quand $\frac{t}{\sigma} + \frac{u\sqrt{-A}}{\sigma}$ est le carré d'une expression semblable. Quand s a des valeurs plus grandes, il dépend des valeurs de h et k ; la s -ième puissance d'une expression semblable doit être $\frac{t}{\sigma} + \frac{u\sqrt{-A}}{s\sigma}$, pour que cette divisibilité ait lieu pareillement. A cela suffisent en tout cas les carrés de ces puissances qui sont $= \frac{T}{\sigma} + \frac{U}{\sigma} \sqrt{-A}$ quand $T^2 + AU^2 = \pm \sigma^2$.

Pour les points de croisement discutés, j'emploierai désormais, quand leur nombre sera plus grand que 1, et que $A_1 = F(\lambda, \mu, \nu)$ sera la valeur spéciale de A , l'expression : *ils sont situés sur le territoire de $A = A_1$ ou $A = F(\lambda, \mu, \nu)$. Le point lui-même $h = k = 0$, auquel correspond une position de $(a_1) = [f(\alpha, \alpha_1, \alpha_2)]$, perpendiculaire sur (A_1) , je l'appellerai le point $\frac{a_1}{A_1}$, en quoi, si besoin est, on peut remplacer a_1 et A_1 par $f(\alpha, \alpha_1, \alpha_2)$ et $F(\lambda, \mu, \nu)$. De même que les points $\frac{f(1, \alpha_1, \alpha_2)}{A}$ correspondent aux points extrêmes des droites $\alpha_1(b) + \alpha_2(c)$ partant du même point initial situé dans le plan (A) , de même ce qui correspond aux lignes droites (b) , (c) ou en général $p(b) + q(c)$ ou $[(b, g, c)(p, q)]$, qui réunissent ces points, et pour lesquels relativement α_2, α_1 , ou généralement $p\alpha_2 - q\alpha_1$, sont constants, ce sont, sur la nappe hyperboloïdale, les lignes réunissant les points $\frac{f(1, \alpha_1, \alpha_2)}{A}$, sur lesquelles, relativement, $\frac{g}{A}, \frac{h}{A}, \frac{p\frac{g}{A} - q\frac{h}{A}}{A}$ sont constants, ou, exprimé par les coefficients de formes appropriées $\mathcal{F}'$, les quantités $\mathcal{H}', \mathcal{A}', p\mathcal{H}' - q\mathcal{A}'$ sont égales à zéro; ces lignes ne se coupent qu'en un seul point ou ne se coupent point du tout, comme les droites correspondantes. Si, dans le cas $A > 0$, aucun des points $\frac{f(1, \alpha_1, \alpha_2)}{A}$ n'est plus situé d'un côté d'une pareille ligne, tandis qu'elle-même passe par deux ou un plus grand nombre de ces points, je l'appelle une *frontière du territoire de A* ; les deux extrêmes des points situés sur elle, je les nomme *sommets d'angles du territoire*; les points qui ne sont situés sur aucune frontière, je les dis *situés**

dans l'intérieur du territoire. Parmi les lignes $\frac{p^4 - q^4}{2} = \text{const.}$, que nous venons d'étudier, se trouvent des lignes limites, précédemment étudiées, de champs dans lesquels sont réduites des différentes formes $\mathcal{F}, \mathcal{F}', \dots$, savoir, des lignes limites, le long desquelles ($\mathcal{A}$) est perpendiculaire sur un autre plan, et qui sont coupées, sur deux points au moins, par d'autres lignes pareilles. Mais ces lignes-là, particulièrement aussi les frontières des territoires de coefficients A positifs, ne se trouvent pas nécessairement parmi ces lignes limites-là, dont seulement trois à la fois passent par les points de croisement à étudier, et qui, si toutefois elles coïncident avec celles-là, peuvent aussi se terminer à des points de fissure, avant d'atteindre ces points de croisement. Dans l'exemple de la *fig. 1*, toutes les frontières, à l'exception d'une, des territoires de coefficients positifs A , coïncident avec des lignes limites de la catégorie précitée, et on les a inscrites comme lignes serpentantes. Seule, la ligne limite avec laquelle, si l'on désigne $\begin{pmatrix} 2, & -1, & -9 \\ 1, & -5, & -1 \end{pmatrix}$ par f , devrait coïncider la frontière du territoire de $A = 8 = F(1, 0, 0)$, conduisant du point $\frac{2}{8}$, à désigner par $\frac{f(1, 0, 0)}{A}$, au point $\frac{3}{8}$, à désigner par $\frac{f(1, 0, -1)}{A}$, part bien du premier point, mais n'atteint pas l'autre; j'ai marqué, en conséquence, une ligne limite existante, qui relie le point $\frac{2}{8}$ au point $\frac{6}{8}$, le second situé sur la frontière suivante, attendu que je tenais simplement à désigner d'une manière quelconque la connexion des points qui appartiennent à des territoires de coefficients positifs A , et que je ne voulais plus introduire aucune ligne étrangère aux deux modes de division des champs. Si l'on observe que, dans la *fig. 1*, les six lignes limites externes sont des axes de symétrie bipartite et leurs six points d'intersection des centres de symétrie quadripartite; si l'on s'imagine, d'après cela, la figure agrandie, on obtient une vue d'ensemble de toutes les frontières des territoires occurrents de coefficients positifs A , savoir, abstraction faite de ceux qui se bornent à des lignes et de celui qui a déjà été mentionné, des territoires de $A = 12$, $A = 8 = \begin{pmatrix} 8, & -12, & -15 \\ 0, & 0, & 12 \end{pmatrix} (1, 0, 0)$ et $A = 5$.

Tout ce qui a été dit ici, pour A et a , F et f , a son analogie complète pour a et A , f et F . C'est aussi sur quoi je fonde la désignation analogue d'un territoire de $a = a_1$ ou $a = f(\alpha, \alpha_1, \alpha_2)$. Ce territoire embrasse, sur la nappe hyperboloïdale, les points de croisement $\frac{a_1}{F_1(1, \mu, \nu)}$ pour toutes les valeurs possibles de μ, ν , savoir les valeurs telles que, sur le plan $[f_1(1, \mu, \nu)]$, la droite (a_1) puisse être perpendiculaire ou qu'on ait $1 \leq a_1 F_1(1, \mu, \nu) \leq I$. Dans le cas $a > 0$, ces points, pour lesquels ici l'expression indiquée aussi ne doit être employée que si leur nombre est plus grand que l'unité, sont situés dans l'intérieur de l'ellipse $2\xi^2 = a + \sqrt{aI}$, et leur nombre est fini; dans le cas $a < 0$, ils sont situés entre deux branches d'hyperboles $\xi = \pm \sqrt{\frac{a + \sqrt{-aI}}{2}}$, et leur nombre est infini.

Sur les lignes qui relient les points pris un à un, $\frac{rh - sk}{a}$ est constant, $rh' - sk' = 0$. Ces lignes, particulièrement les frontières des territoires de coefficients positifs a , coïncident éventuellement, mais non nécessairement, avec les lignes limites des champs, dans lesquels différentes formes $f, f', \dots$ sont réduites. Dans la *fig. 1*, toutes ces frontières, à l'exception de deux, coïncident avec de pareilles lignes et ont été différenciées d'avec les autres lignes limites par des traits renforcés. Mais, dans le territoire de $a = 2$, il est vrai qu'un fragment de la frontière, qui relie le sommet d'angle $\frac{2}{8} = \frac{2}{(8, -40, -3) (1, 0, 0)}$

au sommet d'angle $\frac{2}{8} = \frac{2}{(8, -40, -3) (1, -1, 0)}$ situé sur le com-

plément à ajouter en haut à gauche, coïncide avec les lignes limites $h = 0$ des champs des formes $\begin{pmatrix} 2, -1, -12 \\ 2, -4, -1 \end{pmatrix}$ qui se trouvent dans le dessin et dans le complément précité; mais cette ligne se termine, des deux côtés, à des points de fissure avant d'atteindre les sommets d'angle précités. J'ai donc mieux aimé, d'une manière analogue, comme dans le cas exceptionnel susdit, renforcer la ligne à désigner par $\frac{h}{a} = -1$, par rapport à la forme employée $\begin{pmatrix} 2, -1, -12 \\ 2, -4, -1 \end{pmatrix}$.

ligne qui coïncide avec une ligne limite et qui relie les points situés sur les frontières qui suivent, savoir, les points

$$\frac{2}{17} = \frac{2}{\begin{pmatrix} 8, & -40, & -3 \\ 0, & -6, & -20 \end{pmatrix} (1, 0, -1)} \quad \text{et} \quad \frac{2}{17} = \frac{2}{\begin{pmatrix} 8, & -40, & -3 \\ 0, & -6, & -20 \end{pmatrix} (1, -1, -1)}$$

Dans le territoire de $a = 3 = \begin{pmatrix} 3, & -4, & -5 \\ 0, & 0, & 0 \end{pmatrix} (1, 0, 0)$, la frontière $\frac{h+k}{a} = -1$, qui relie les points $\frac{3}{5}$ et $\frac{3}{8}$, ne coïncide avec aucune ligne limite; j'ai renforcé, à sa place, une autre ligne limite qui relie ces points. Outre les territoires qui se bornent à des lignes et ceux qui ont déjà été mentionnés, nous trouvons (fig. 1) encore un territoire d'un coefficient positif a , savoir, celui de

$$a = 3 = \begin{pmatrix} 3, & -1, & -6 \\ 1, & -3, & -2 \end{pmatrix} (1, 0, 0).$$

Maintenant, non plus par la division de toute la nappe hyperboloïdale en nombreux petits champs, mais immédiatement, je cherche l'existence et la situation respective des territoires moins nombreux des coefficients positifs A et a .

Comme la substitution $\begin{vmatrix} 1 & -\alpha_1 & -\alpha_2 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{vmatrix}$ est adjointe à la substitu-

tion $\begin{vmatrix} 1 & 0 & 0 \\ \alpha_1 & 1 & 0 \\ \alpha_2 & 0 & 1 \end{vmatrix}$, on obtient $F(-\alpha_1, 1, 0) < 0$, et $F(-\alpha_2, 0, 1) < 0$,

quand on a $f(1, \alpha_1, \alpha_2) > 0$; par conséquent, si $A > 0$, nous avons $F(1, \mp 1, 0) < 0$ pour le signe supérieur, ou pour le signe inférieur, ou pour tous les deux, suivant qu'il existe une valeur entière positive ou une négative, ou une positive et une négative de α , qui, jointe à une valeur quelconque, même fractionnaire, de α_2 , rend $f(1, \alpha_1, \alpha_2)$ positif; de même on a $F(1, 0, \mp 1) < 0$, suivant qu'il existe une valeur positive ou une négative, ou une valeur positive et une négative de α_2 , laquelle, jointe à une valeur quelconque, même fractionnaire de α_1 , rend positif $f(1, \alpha_1, \alpha_2)$.

Si le point $\frac{a}{A}$ est situé dans l'intérieur du territoire de A , chaque

fois le troisième de ces six cas a lieu; par conséquent, pour tous les signes, $F(1, \mp 1, 0)$ et $F(1, 0, \mp 1)$ sont négatifs, et attendu qu'il en est ensuite de même pour chaque forme équiva-

lente F_1 , résultant de F par une substitution quelconque $\begin{vmatrix} 1 & 0 & 0 \\ 0 & \mu_1 & \mu_2 \\ 0 & \nu_1 & \nu_2 \end{vmatrix}$,

il ne peut pas alors y avoir de nombres premiers entre eux, ni par conséquent de nombres entiers quelconques μ, ν différents de zéro qui rendraient positif $F(1, \mu, \nu)$; si donc le point $\frac{a}{A}$ est situé dans l'intérieur d'un territoire de A , il n'appartient à aucun territoire de a ; on trouvera de même que, s'il est situé dans l'intérieur d'un territoire de a , il n'appartient à aucun territoire de A .

Si, au contraire, le point $\frac{a}{A}$ est situé sur une frontière du territoire

de A , il y a des substitutions $\begin{vmatrix} 1 & 0 & 0 \\ 0 & \beta_1 & \gamma_1 \\ 0 & \beta_2 & \gamma_2 \end{vmatrix}$, au moyen desquelles f se

change en une forme f_1 , de sorte que, si l'on a $\alpha_1 > 0$ et $f_1(1, \pm 1, 0) > 0$ pour un des deux signes au moins, il n'y a aucune valeur entière positive de α_2 , ou il n'y a aucune valeur entière négative de α_2 différant toujours de zéro, qui, jointe à une valeur entière quelconque de α_1 , rende $f_1(1, \alpha_1, \alpha_2)$ positif. Je fais provisoirement une supposition encore plus étroite: c'est qu'il n'y ait aucune valeur entière positive de α_2 différente de zéro, ou qu'il n'y ait aucune valeur entière négative de α_2 , différente de zéro, qui, jointe à une valeur réelle quelconque, entière ou fractionnaire, de α_1 , rende $f_1(1, \alpha_1, \alpha_2)$ positif, ou, ce qui revient évidemment au même, c'est qu'avec un des deux signes $f_1(1, \alpha_1, \pm 1)$ ne devienne positif pour aucune valeur réelle, entière ou fractionnaire, de α_1 . Pour abréger, j'appellerai désormais *supposition* $\mathfrak{U}$ cette supposition faite pour une frontière d'un territoire de A , et pareillement la supposition analogue faite pour une frontière d'un territoire de a . De celle-là, savoir, de l'absence d'une racine réelle α_1 de l'équation $f_1(1, \alpha_1, \pm 1) = 0$, provient maintenant $F_1(1, 0, \mp 1) > 0$, et vice versa, et l'analogie existe après la permutation de f et F .

Par l'application de la même considération aux points $\frac{f_1(1, u, 0)}{A_1}$,

qui doivent représenter tous les points d'une frontière du territoire de A_i répondant à l'équation $\mathfrak{H}_i = 0$, tandis que u admet une série de u' valeurs entières consécutives renfermant aussi la valeur zéro, on arrive à la conclusion suivante. Si la supposition $\mathfrak{D}$ est réalisée pour cette frontière, tous ces points appartiennent aussi à des territoires des différents coefficients $f_i(1, u, 0)$, dans lesquels sont encore situés les points $\frac{f_i(1, u, 0)}{F_i(1, 0, \mp 1)}$, reliés respectivement par les lignes $k + b, u = 0$ avec les points $\frac{f_i(1, u, 0)}{A_i}$, et entre eux par la ligne $\mathfrak{H}_i \mp \mathfrak{C}_i = 0$, lesquels points évidemment appartiennent tous à un nouveau territoire de $F_i(1, 0, \mp 1)$. Si $F_i(1, 0, -v)$ demeure positif pour une série de v' valeurs entières v , se suivant et renfermant aussi la valeur zéro, il y aura évidemment $u' \cdot v'$ points $\frac{f_i(1, u, 0)}{F_i(1, 0, -v)}$, dont chacun appartient à un des u' territoires des coefficients $f_i(1, u, 0)$ et à un des v' territoires des coefficients $F_i(1, 0, -v)$.

De $f_i(1, \pm 1, 0) > 0$ suit non-seulement $F_i(1, \mp 1, 0) < 0$, mais, la substitution $\begin{vmatrix} 1 & 0 & -v \\ \pm 1 & 1 & 0 \\ 0 & 0 & 1 \end{vmatrix}$ ayant pour adjointe $\begin{vmatrix} 1 & \mp 1 & 0 \\ 0 & 1 & 0 \\ v & \mp v & 1 \end{vmatrix}$, nous aurons encore $F_i(1, \mp 1, v) < 0$ pour toute valeur réelle, entière ou fractionnaire de v . Les points $\frac{a_i}{F_i(1, 0, -v)}$ sont donc situés sur une frontière du territoire de a_i , ou tout ce territoire se borne à cette frontière, suivant que la condition $f_i(1, \pm 1, 0) > 0$ est réalisée pour un seul des signes ou pour tous les deux à la fois. Car, si $F_i(1, \mp \mu, v)$ était positif, pour une valeur entière positive quelconque μ jointe à une valeur entière quelconque v , il faudrait aussi que $F_i(1, \mp 1, v)$ fût > 0 , du moins pour quelque valeur fractionnaire v . On reconnaît, d'après cela, que, pour les deux valeurs extrêmes de u , les points $\frac{f_i(1, u, 0)}{F_i(1, 0, -v)}$ sont situés sur des frontières des territoires des coefficients positifs $f_i(1, u, 0)$. On voit aussi que les territoires des coefficients positifs $f_i(1, u, 0)$ se bornent tous à des lignes, pour toutes les valeurs autres que les deux valeurs extrêmes de u . Enfin, par analogie, on conclura que, pour les deux valeurs extrêmes de v ,

les points $\frac{f_1(1, u, 0)}{F_1(1, 0, -v)}$ sont tous situés sur des frontières des territoires des coefficients positifs $F_1(1, 0, -v)$, et que les territoires des coefficients positifs $F_1(1, 0, -v)$ se bornent à des lignes pour toutes les valeurs autres que les deux valeurs extrêmes de v .

Supposons que le territoire de A_1 , considéré en premier lieu, ne se réduise pas à une ligne, et soit zéro une des valeurs extrêmes de u .

Alors du point $\frac{a_1}{A_1}$ part une seconde frontière du territoire de A_1 , sur laquelle est situé un point $\frac{f_1(1, \beta_1, \beta_2)}{A_1}$, β_1 et β_2 étant certains nombres premiers entre eux, lequel point peut aussi être désigné par $\frac{f_2(1, 1, 0)}{A_2}$,

si la forme f_2 provient de la forme f_1 par une substitution $\begin{vmatrix} 1 & 0 & 0 \\ 0 & \beta_1 & \gamma_1 \\ 0 & \beta_2 & \gamma_2 \end{vmatrix}$.

Si la supposition $\mathfrak{U}$ ne cesse pas d'être remplie pour cette nouvelle frontière, il y aura aussi un point $\frac{a_2}{F_2(1, 0, \mp 1)}$ qui, appartenant au territoire de $a = a_1 = a_2$, n'est pas situé sur la frontière considérée ci-dessus de ce territoire. Donc aussi ce territoire ne se réduit pas à une ligne, ni donc manifestement aucun des quatre territoires de $a = f_1(1, u, 0)$ et $A = F_1(1, 0, -v)$ pour les valeurs extrêmes de u et v en tant que la supposition $\mathfrak{U}$ est remplie. Les quatre frontières considérées de ces quatre territoires forment donc un quadrilatère, tel que ces territoires eux-mêmes, considérés comme des surfaces entourées de leurs frontières, sont situés en dehors du quadrilatère. Si $u' > 2$, il y a, dans l'intérieur du quadrilatère, $u' - 2$ territoires de coefficients $f_1(1, u, 0)$. Ces territoires se bornent à des lignes, qui ne peuvent ni se couper l'une l'autre ni couper les frontières des deux autres territoires de coefficients $f_1(1, u, 0)$, comme le prouvent les équations de ces lignes $k_1 + b_1 u = 0$. Il en est de même, par analogie, quand $v' > 2$, des lignes $\mathfrak{H}_1 - \mathfrak{C}, v = 0$, qui coupent les premières en forme de grille.

Dans l'exemple de la fig. 1, on n'a nulle part en même temps $u' > 2$ et $v' > 2$: ainsi ces deux sortes de lignes ne se trouvent jamais simultanément. Pour la détermination de territoires ultérieurs, ceux qui sont situés dans l'intérieur d'un pareil quadrilatère et se réduisent

à des lignes n'ont aucune importance; aussi n'en tiendrai-je plus aucun compte dans les exemples qui viendront plus tard.

Si l'on applique les observations faites aux deux frontières qui se rencontrent aux sommets d'angles de territoires de coefficients positifs A et a , on reconnaît, *en tant que la supposition 3 se réalise, ce qui suit : chaque sommet d'angle d'un territoire de A est en même temps un sommet d'angle d'un territoire de a et vice versa. De plus, quand deux frontières dirigées l'une vers l'autre, de territoires de deux coefficients a , partent des deux sommets d'angle d'une frontière d'un territoire de A , elles se terminent aussi en deux sommets d'angles d'une frontière, dirigée vers la précédente, d'un territoire d'un autre coefficient A .* Dans la fig. 1, la première proposition se réalise de tous points, mais pas la deuxième. En effet, la supposition 3 ne se réalise pas à la frontière précitée, entre les points $\frac{2}{8}$ et $\frac{3}{8}$ qui, en posant $f_1 = \begin{pmatrix} 2, & -9, & -1 \\ 1, & -1, & -5 \end{pmatrix}$, peuvent être désignés par $\frac{f_1(1, 0, 0)}{A_1}$ et $\frac{f_1(1, -1, 0)}{A_1}$; car nous avons $F_1(1, 0, -1) < 0$, tandis que $f_1(1, \alpha_1, \alpha_2)$ ne devient positif pour aucune valeur entière positive de α_2 , jointe à une valeur entière quelconque de α_1 .

Je regarde comme normal le cas où la supposition 3 se réalise; le cas contraire est pour moi l'exception. En effet, si elle ne se réalise pas dans un territoire de A pour une frontière sur laquelle sont situés des points $\frac{f(1, u, 0)}{A}$, tandis que nous avons $f(1, u, 1) \leq -1$ pour toutes les valeurs entières de u , il faut que la plus grande valeur, dont $f(1, u, 1)$ est susceptible avec la variabilité continue de u , soit positive; on a donc $f\left(1, \frac{g+k}{-b}, 1\right) = a + 2h + c - \frac{(g+k)^2}{b} \geq \frac{1}{-b}$; par contre, attendu que déjà pour la valeur entière la plus rapprochée, $u_1 = \frac{g+k}{-b} \pm \delta$, de u , on a $f(1, u_1, 1) \leq -1$, on obtient pour la même expression $a + 2h + c - \frac{(g+k)^2}{b} \leq -b\delta^2 - 1$, tandis que $\delta^2 \leq \frac{1}{4}$, et que pour $-b$ des limites sont posées par la condition $f(1, u, 0) \geq 1$ à prendre pour les valeurs extrêmes de u .

Quand la supposition 3 se réalise pour une frontière, le quadrilatère dont nous avons parlé existe et la supposition est par conséquent

réalisée pour les quatre frontières qui le forment. *En tant que la supposition $\mathfrak{U}$ est valable, la surface de l'hyperboloïde se partage donc en trois genres de parties, savoir d'une part en territoires de A , entourés de certaines frontières, de l'autre en territoires de a et enfin en les quadrilatères précités.* Ces fragments de surface s'excluent tous réciproquement et remplissent néanmoins sans lacune toute la surface. Si l'on a développé les territoires à une distance suffisante dans toutes les directions, des territoires déjà vus doivent se répéter et le développement ultérieur peut se ramener au développement antérieur. Il n'est pas nécessaire de changer les formes f , donnant des points de croisement $\frac{a}{A}$ de telle sorte que les formes f soient réduites; car, leur observation ayant démontré la possibilité de ces développements existant sous la supposition $\mathfrak{U}$, on n'a plus du tout besoin d'en tenir compte.

A propos des points de croisement $\frac{a'}{A}$, dans lesquels a' et A' concordent numériquement avec les coefficients a et A qui se sont présentés auprès d'un point de croisement antérieur $\frac{a}{A}$, je veux aussi décider immédiatement sur l'accord numérique des formes f' et f qui me servent de base; et dans ce but je transforme, ce qui est utile aussi dans le développement de tous les territoires de coefficients A , toutes les formes qui se

présentent, au moyen de substitutions $\begin{vmatrix} 1 & 0 & 0 \\ 0 & \beta_1 & \gamma_1 \\ 0 & \beta_2 & \gamma_2 \end{vmatrix}$, de telle sorte que

$-(h, g, c)$ soit réduit. Pour le développement des territoires de coefficients a , une transformation passagère, qui réduit $-(B, G, C)$, est souvent utile. Il n'est pas nécessaire non plus, en rencontrant successivement chaque frontière d'un territoire de A , de changer, au moyen d'une substitution spéciale, la forme f en la forme f_1 adoptée plus haut. Car, parce qu'on obtient les expressions $f_1(1, u, 0)$ et $F_1(1, 0, -v)$, correspondant aux u, v points de croisement considérés, au moyen

des substitutions $\begin{vmatrix} 1 & 0 & 0 \\ u & 1 & 0 \\ 0 & 0 & 1 \end{vmatrix}$ et $\begin{vmatrix} 1 & 0 & v \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{vmatrix}$ appliquées en ordre quel-

conque à f_1 ou des substitutions $\begin{vmatrix} 1 & -u & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{vmatrix}$ et $\begin{vmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ -v & 0 & 1 \end{vmatrix}$ appliquées

à F_1 , ces expressions, après insertion d'une substitution $\begin{vmatrix} 1 & 0 & 0 \\ 0 & \beta_1 & \gamma_1 \\ 0 & \beta_2 & \gamma_2 \end{vmatrix}$,

deviennent $f(1, u\beta_1, u\beta_2)$ et $F(1, v\beta_2, -v\beta_1)$.

Je veux développer quelques exemples, dans lesquels, ce qui est aussi le cas le plus fréquent lorsqu'il s'agit d'un petit nombre de territoires numériquement différents, la supposition $\mathfrak{U}$ se réalise de tous points. Dans les figures, les lignes serpentantes, dont quelques-unes sont rattachées entre elles à des angles et entourent un nombre positif A , désignent les frontières du territoire de ce coefficient A ; il en est de même des lignes pleines, des frontières des territoires des coefficients positifs a ; enfin les lignes ponctuées désignent des axes de symétrie. Sur les frontières, j'ai fait ressortir particulièrement les points $\frac{f(1, u, 0)}{A_1}$ et $\frac{a_1}{F_1(1, 0, -v)}$; de même sur les axes de symétrie les points

$\frac{a}{A}$, qui y sont situés. Dans la représentation, je pars immédiatement d'une des formes de la classe choisie, d'où naissent le plus simplement les symétries qui se produisent; je pars donc, dans l'exemple de la fig. 2, de la forme $f = \begin{pmatrix} 1, -2, -2 \\ 1, 0, 1 \end{pmatrix}$. Le point $\frac{1}{3} = \frac{f(1, 0, 0)}{A}$ appartient, avec les points $\frac{1}{3} = \frac{f(1, 1, 0)}{A}$ et $\frac{1}{3} = \frac{f(1, 1, 1)}{A}$, à un territoire de $A = 3$. Comme f se change en elle-même par les substitutions corres-

pondantes $\begin{vmatrix} 1 & 0 & 0 \\ 1 & -1 & 1 \\ 0 & 0 & 1 \end{vmatrix}$ et $\begin{vmatrix} 1 & 0 & 0 \\ 1 & 0 & -1 \\ 1 & -1 & 0 \end{vmatrix}$, ou quand le déterminant doit

être égal à 1 par les substitutions $\begin{vmatrix} -1 & 0 & 0 \\ -1 & 1 & -1 \\ 0 & 0 & -1 \end{vmatrix}$ et $\begin{vmatrix} -1 & 0 & 0 \\ -1 & 0 & 1 \\ -1 & 1 & 0 \end{vmatrix}$, on

reconnait déjà une symétrie tripartite et même une symétrie sexpartite, chacune des trois frontières symétriques étant de nouveau symétrique vers ses deux extrémités. Comme la forme adjointe $F = \begin{pmatrix} 3, -2, -3 \\ -1, 1, 2 \end{pmatrix}$ donne une valeur positive pour $F(1, 0, 1)$ et de même pour $F(1, 1, -1)$, ce qu'on reconnait déjà d'après la symétrie trouvée, la supposition $\mathfrak{U}$ est réalisée pour les deux frontières partant du point

$\frac{f(1, 0, 0)}{F(1, 0, 0)}$. Comme de F , par chacune des deux substitutions $\begin{vmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 1 & 0 & 1 \end{vmatrix}$

et $\begin{vmatrix} 1 & 0 & 0 \\ 1 & 1 & 1 \\ -1 & 0 & -1 \end{vmatrix}$, provient la forme $\begin{pmatrix} 2, & -2, & -3 \\ -1, & -2, & 1 \end{pmatrix}$, qui se change en elle-même par la substitution $\begin{vmatrix} 1 & 0 & 0 \\ 1 & -1 & -1 \\ 0 & 0 & 1 \end{vmatrix}$, on reconnaît que le territoire de $a = 1 = f(1, 0, 0)$, outre l'axe de symétrie passant par le point $\frac{1}{F(1, 0, 0)}$, est encore coupé par un autre axe semblable qui coupe le premier dans un centre de symétrie quadripartite, au point de croisement $\frac{1}{F(1, 1, 0)}$. Ce territoire a donc six angles, savoir : dans la série où sont reliés les uns aux autres, par des frontières, les points $\frac{1}{F(1, 0, 0)}$, $\frac{1}{F(1, 0, 1)}$, $\frac{1}{F(1, 1, 1)}$, $\frac{1}{F(1, 2, 0)}$, $\frac{1}{F(1, 2, -1)}$, $\frac{1}{F(1, 1, -1)}$. Des quadrilatères qui se rattachent aux frontières du territoire de $A = 3$, on connaît déjà les quatre angles, d'après les symétries connues, par exemple, outre les deux sommets d'angles $\frac{f(1, 0, 0)}{F(1, 0, 0)}$ et $\frac{f(1, 1, 0)}{F(1, 0, 0)}$, les sommets d'angles symétriques de même l'un à l'autre $\frac{f(1, 0, 0)}{F(1, 0, 1)}$ et $\frac{f(1, 1, 0)}{F(1, 0, 1)}$. On reconnaît que la supposition $\mathfrak{M}$ est aussi réalisée par les frontières, n'appartenant pas à ce quadrilatère et sortant de ces nouveaux sommets d'angles. C'est ce que l'on voit en considérant, par exemple, la forme $\begin{pmatrix} 1, & -2, & -1 \\ 0, & -1, & 1 \end{pmatrix}$, provenant de f par la substitution $\begin{vmatrix} 1 & 0 & -1 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{vmatrix}$, et qui, ainsi que nous l'avons déjà vu, mène par la substitution ultérieure $\begin{vmatrix} 1 & -1 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{vmatrix}$ au sommet d'angle suivant $\frac{1}{F(1, 1, 1)}$ du territoire de $a = f(1, 0, 0)$. Cette forme donne aussi par la substitution $\begin{vmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ -1 & 0 & 1 \end{vmatrix}$ une forme à premier coefficient positif. Cette nouvelle forme $\begin{pmatrix} 2, & -2, & -1 \\ 0, & 0, & 1 \end{pmatrix}$ donne un nouvel axe de symétrie, et l'on est ainsi parvenu au terme des considérations numériques nécessaires. Comme lors de la coïncidence de deux formes, les champs, les ter-

ritoires, les frontières et les angles doivent coïncider, on reconnaît ce que la *fig. 2* représente. Les six points d'intersection des six axes indiqués, de symétrie, sont des centres de symétrie quadripartite, ce qui doit être indiqué par les angles droits choisis, comme dans d'autres exemples, une symétrie de λ parties est indiquée par les angles de $\frac{4}{\lambda}$ droits. Si l'on se figure le dessin tellement ployé que deux axes pareils, perpendiculaires l'un à l'autre, deviennent des lignes droites, on n'aura pas de peine à s'imaginer la quadruplication de la figure, et l'on peut effectuer la même chose toutes les fois que deux axes se coupent et sont situés à la limite de la figure agrandie. Je n'ai pas voulu faire usage de la symétrie sextuple, qui a lieu autour du centre de la figure, pour la réduire, afin de faciliter la vue d'ensemble. Pour deux autres exemples, les *fig. 3* et *4* pourront suffire seules; j'y ajoute aux points externes $\frac{a}{A}$, leur dénomination ou les formes $\begin{pmatrix} a, b, c \\ g, h, k \end{pmatrix}$ donnant $h = k = 0$. L'indication de A , et α , n'est alors plus nécessaire.

Quant aux lieux où la supposition $\mathfrak{V}$ n'est pas réalisée, je pourrais renvoyer simplement à la méthode primitive de la division en champs; mais, même en passant par-dessus ceux-ci, on peut continuer le développement suivant le mode abrégé et purement arithmétique employé en dernier lieu. Cependant, sous cette réserve, je vais me résumer et ne poursuivrai pas les cas d'exception les plus rares. Si α et A sont positifs et en admettant que le point $\frac{\alpha}{A}$ soit un sommet d'angle du territoire de A , d'où sort sa frontière contenant les points $\frac{f(1, u, 0)}{A}$ pour $0 \leq u \leq u'$, tandis qu'il n'y a pas, dans le territoire, de points $\frac{f(1, \alpha_1, \alpha_2)}{A}$, à valeur négative de α_2 , jointe à une valeur quelconque de α_1 ; en admettant de plus que $F(1, 0, 1) < 0$, de sorte que la supposition $\mathfrak{V}$ ne se réalise pas pour cette frontière, le quadrilatère précité, devant se rattacher à la frontière du côté extérieur, n'existe pas; alors la supposition $\mathfrak{V}$ n'est pas réalisée non plus pour une seconde frontière dirigée vers la première et appartenant au territoire de α , qui, si toutefois il existe, doit avoir $\frac{\alpha}{A}$ comme sommet d'angle. Il en est de même d'une troisième

frontière dirigée vers la seconde et appartenant au territoire de A' , qui, s'il existe, doit avoir pour sommet d'angle le point final de la seconde, et ainsi de suite. Évidemment l'ensemble de toutes ces frontières consécutives de territoires alternatifs de différents coefficients A et a donne une limite continue au delà de laquelle, dans une direction, la supposition $\mathfrak{H}$ ne se réalise pas. Un territoire que l'on étudie peut maintenant aussi dégénérer en une ligne qui doit être regardée comme sa frontière vers ses deux côtés opposés. Il se peut alors que, seulement vers un côté ou, comme dans l'exemple de la *fig. 5*, pour le territoire $A = 12$, vers les deux côtés d'une telle ligne, la supposition $\mathfrak{H}$ ne se réalise pas. Dans le dernier cas, cette ligne constitue, vers les deux côtés, une limite de l'espèce que nous avons précédemment indiquée. On trouve encore une pareille limite, si aucun territoire ne se rattache au point extrême d'une des frontières étudiées; savoir si, comme, par exemple, pour la forme $\begin{pmatrix} 44, & -45, & -46 \\ -22, & 2, & 1 \end{pmatrix}$ avec l'adjointe $\begin{pmatrix} 1586, & -2028, & -1981 \\ 970, & 68, & 2 \end{pmatrix}$, ce point extrême $\frac{a}{A}$ appartient à un territoire de A , et que la condition $F(1, \mu, \nu) > 0$ ne se réalise que pour $\mu = \nu = 0$, on poursuit la frontière suivante partant de ce point, du territoire de A . Avec cette limite continue le cas peut d'abord se présenter où elle rentre en elle-même, de manière que les territoires, auxquels appartiennent les frontières qui la constituent, soient situés en dehors de la ligne fermée. Ce cas se présente dans l'exemple de la *fig. 1*, où les frontières respectives dont l'ensemble ne se montre qu'après le complément à effectuer en haut, du côté gauche, et dont une a déjà été mentionnée, constituent un hexagone formé alternativement de frontières de territoires de coefficients positifs A et a . Un cas analogue se présente dans l'exemple tracé *fig. 5*, qui n'a pas besoin d'être expliqué davantage, et dans la *fig. 6*, où toutefois un côté de l'hexagone, une frontière d'un territoire de $a = 0$, tombe à une distance infinie. Il est facile de décider, dans les hypothèses admises ci-dessus, laquelle des deux frontières du territoire de a , issues du point $\frac{a}{A}$, et qui contiennent les points $\frac{a}{F(1, \mu, \nu)}$, est dirigée vers la frontière contenant les points $\frac{f(1, u, 0)}{A}$ du territoire de A . Car le rapport $\frac{\mu}{\nu}$ des deux nom-

bres nécessairement positifs μ, ν , qui serait égal à zéro pour la frontière cherchée, si la supposition $\mathfrak{M}$ était remplie, doit être le plus petit possible, si la supposition $\mathfrak{M}$ n'est pas remplie. Comme le fragment de surface enfermé dans la limite rentrant en elle-même, hexagone dans les exemples cités, peut aussi bien être négligé que le quadrilatère précité, dans le cas de la supposition $\mathfrak{M}$, on voit que le cas étudié ne présente aucun obstacle dans le développement ultérieur des territoires. Deuxièmement, la limite mentionnée peut maintenant aussi s'étendre des deux côtés, à l'infini, de sorte que des formes concordant numériquement reviennent auprès de la limite, d'une manière périodique. Pour ce cas, j'allègue les formes $\begin{pmatrix} 2r, & -2r, & 2r-s \\ 0, & 0, & r \end{pmatrix}$, dans lesquelles, si $4r < s < \frac{9}{2}r$, il n'y a que

les substitutions alternatives $\begin{vmatrix} 1 & 0 & 0 \\ 1 & 1 & 0 \\ 0 & 0 & 1 \end{vmatrix}$ et $\begin{vmatrix} 1 & 1 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{vmatrix}$ qui puissent

alternativement amener de nouvelles valeurs positives à la place de a et A . Troisièmement, cette limite peut rentrer en elle-même, de telle sorte qu'elle enferme les territoires, pour lesquels la supposition $\mathfrak{M}$ est réalisée. Il est possible encore qu'on n'obtienne qu'un seul et unique territoire, se réduisant même à une ligne; enfin, il peut y avoir même des points détachés $\frac{a}{A}$ de nature à faire que $f(1, \alpha_1, \alpha_2)$ ne soit positif que pour $\alpha_1 = \alpha_2 = 0$, et que $F(1, \mu, \nu)$ ne soit positif que pour $\mu = \nu = 0$, de quoi la forme $\begin{pmatrix} 44, & -47, & -47 \\ 20, & 1, & 1 \end{pmatrix}$ avec l'adjointe $\begin{pmatrix} 1809, & -2069, & -2069 \\ -879, & 67, & 67 \end{pmatrix}$ peut servir d'exemple.

Si $\frac{f(1, u, 0)}{A}$ sont les points d'une frontière d'un territoire de A , auquel n'appartiennent des points $\frac{f(1, \alpha_1, \alpha_2)}{A}$ pour aucune valeur négative de α_2 , jointe à une valeur quelconque de α_1 , et si la supposition $\mathfrak{M}$ ne se réalise pas pour celle-ci, si donc on a $F(1, 0, 1) < 0$, on pourra continuer le développement au delà de cette frontière en appliquant à F une des substitutions à trouver selon II, b

$$\begin{vmatrix} \lambda & 0 & \lambda_2 \\ 0 & 1 & 0 \\ \nu & 0 & \nu_2 \end{vmatrix} \text{ au déterminant } 1, \text{ qui donnent } F(\lambda, 0, \nu) \geq 1, F(\lambda_2, 0, \nu_2) \leq -1,$$

savoir celle qui possède les plus petits coefficients, et fait $f(\nu_2, \alpha_1, -\lambda_2)$

positif pour quelque valeur entière de α_1 . J'appelle F' la forme qui provient ainsi de F . Si, pour arriver au but projeté, qui sera, le plus souvent, atteint dès le premier pas, des coefficients plus petits λ , ν , λ_2 , ν_2 ne suffisaient pas, en tous cas les plus petits qu'on pourra trouver suffiront, d'après ce qu'on a vu plus haut pour les territoires des

coefficients négatifs A , et qui, à l'aide de la substitution $\begin{vmatrix} \nu_1 & 0 & -\nu \\ \alpha_1 & 1 & \gamma_1 \\ -\lambda_2 & 0 & \lambda \end{vmatrix}$,

changent la forme $\begin{pmatrix} a, b, c \\ g, h, k \end{pmatrix}$ en la forme $\begin{pmatrix} a, b, c \\ \epsilon g, h, \epsilon k \end{pmatrix}$. Je passe sur la diminution qu'on peut obtenir lorsque le signe inférieur est possible.

La série des substitutions $\begin{vmatrix} 1 & 0 & 1 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{vmatrix}$ et $\begin{vmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 1 & 0 & 1 \end{vmatrix}$, desquelles se

compose la substitution $\begin{vmatrix} \lambda & 0 & \lambda_2 \\ 0 & 1 & 0 \\ \nu & 0 & \nu_2 \end{vmatrix}$, se termine par la substitution

$\begin{vmatrix} 1 & 0 & 1 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{vmatrix}$; parce que celle-ci peut seule changer les nombres ν_2 , λ_2 qui

apparaissent dans $f(\nu_2, \alpha_1, -\lambda_2)$ et que les conditions pour $F(\lambda, 0, \nu)$ et $F(\lambda_2, 0, \nu_2)$ sont réalisées après chacune de ces substitutions. Il en résulte qu'on obtient $F'(1, 0, -1) < 0$ et $f'(1, u, \nu) < 0$ pour toutes les valeurs positives entières de ν , jointes à des valeurs entières de u ; car, si ν n'est pas plus grand que le nombre p des substitutions

$\begin{vmatrix} 1 & 0 & 1 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{vmatrix}$, qui ont été ajoutées finalement, cela provient de ce que

les conditions auraient d'ailleurs déjà été réalisées par des valeurs plus petites de p ; mais, quand $\nu > p$, la même circonstance résulte de

ce que la substitution $\begin{vmatrix} \lambda & 0 & \lambda_2 \\ 0 & 1 & 0 \\ \nu & 0 & \nu_2 \end{vmatrix} \cdot \begin{vmatrix} 1 & 0 & -\nu \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{vmatrix}$, appliquée à F , donnerait alors un troisième coefficient positif. S'il y a au moins deux points

$\frac{f'(1, u, 0)}{F'(1, 0, 0)}$, ces points appartiennent donc de nouveau à une frontière,

pour laquelle la supposition $\textcircled{3}$ ne se réalise pas; on a donc de nouveau une limite de l'espèce précitée. S'il y avait encore d'autres territoires insérés entre cette frontière et la frontière primitive contenant les

points $\frac{f(1, u, 0)}{A}$, ce qu'on peut décider de la manière la plus simple par l'étude des champs qui apparaissent sur la ligne $\eta = 0$, on aurait déjà pu trouver une limite plus rapprochée, en insérant des substitutions aux endroits convenables, lesquelles substitutions auraient aussi amené d'autres coefficients à la place de b . Je m'abstiens d'examiner plus en détail ce cas, ainsi que celui dans lequel, excepté pour $u = 0$, il n'y a pas de point $\frac{f'(1, u, 0)}{F'(1, 0, 0)}$, lequel cas d'ailleurs, le plus souvent, n'établit pas de différence essentielle avec l'autre, et je renvoie à la méthode générale fondée sur la division en champs.

Soit maintenant que la nouvelle limite, dans tout son parcours, fasse vis-à-vis à l'ancienne, soit que d'autres frontières qui la composent aient pour vis-à-vis des frontières d'autres limites, frontières à trouver de la même manière, on pourra toujours développer les groupes de territoires, séparés les uns des autres comme par des fleuves, mais aussi pouvant être réunis, en cas de besoin, comme par des ponts, assez au loin pour que toutes les formes nouvelles concordent numériquement avec des formes qui s'étaient déjà présentées.

e. — Les transformations des formes en elles-mêmes.

Outre les substitutions déjà étudiées, IV, b , lesquelles, pour une forme réduite f , ne sont possibles que lorsque leur champ coïncide avec lui-même, après une permutation entre a, b, c, d ou avec un champ immédiatement contigu, *il existe encore une infinité de substitutions par lesquelles f se transforme en elle-même, chaque nouveau champ, pour lequel est réduite une forme concordant numériquement avec f , déterminant une nouvelle substitution semblable*, qui correspond en quelque sorte à une ligne reliant le champ primitif au nouveau champ, ou, si au lieu des champs on ne considère que des points de croisement analogues, à une ligne qui relie les deux points de croisement. Comme toute la nappe hyperboloïdale est couverte sans lacune par les figures qui ont été expliquées et dessinées, IV, c et d , et que l'on doit répéter à l'infini, une ligne, reliant un coin d'une pareille figure à sa répétition, peut être formée par les lignes limites de cette figure et leurs répétitions à prendre dans la série et la direction con-

venables. La *substitution* correspondante peut donc être composée de *substitutions de l'espèce étudiée*, IV, *b*, qui correspondent à la permutation de limites analogues de champs ou au franchissement d'axes de symétrie et des *substitutions en nombre limité* P, Q, R, S, ..., qui répondent aux *lignes limites externes de la figure*, et dont l'une au moins, puisque ces lignes limites forment un polygone, peut être exprimée par les autres.

Ici se poserait géométriquement la question de savoir comment cette ligne de communication peut être composée du plus petit nombre des fragments précités, à propos de quoi il faudrait ou bien traiter de la même manière les fragments hétérogènes, ou bien procéder de manière qu'avant tout les fragments d'une même catégorie figurent au plus petit nombre possible, puis de même ceux d'une deuxième catégorie.... A cette question géométrique répond exactement celle de rechercher comment les substitutions à étudier se composent des substitutions P, Q, R, S, etc. Toutefois, comme non-seulement l'espèce des symétries possibles, mais encore le nombre des substitutions P, Q, R, S, ... varient suivant les différentes classes, il n'est pas aisé d'établir d'une manière tout à fait générale l'expression désirable qui donne précisément une fois chacune des substitutions possibles que l'on cherche, tandis qu'il n'y a aucune difficulté à l'établir dans chaque cas donné.

Je veux donc obtenir cette expression seulement pour les exemples

traités jusqu'ici. Soit (*fig. 1*), désignée par P, la substitution $\begin{vmatrix} 1 & 1 & 0 \\ 3 & 4 & 0 \\ 0 & 0 & 1 \end{vmatrix}$

composée des substitutions $\begin{vmatrix} 1 & 0 & 0 \\ 3 & 1 & 0 \\ 0 & 0 & 1 \end{vmatrix}$ et $\begin{vmatrix} 1 & 1 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{vmatrix}$, au moyen de

laquelle la forme $f = \begin{pmatrix} 12, & -1, & -5 \\ 0, & 0, & 0 \end{pmatrix}$ se change en $f_1 = \begin{pmatrix} 3, & -4, & -5 \\ 0, & 0, & 0 \end{pmatrix}$, et qui correspond à la ligne limite placée en bas à droite, laquelle relie le point $\frac{12}{5} = \frac{a}{A}$ au point $\frac{3}{20} = \frac{a_1}{A_1}$. Désignons par Q la substitution

$\begin{vmatrix} 2 & 0 & 5 \\ 0 & 1 & 0 \\ 3 & 0 & 8 \end{vmatrix} = \begin{vmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 1 & 0 & 1 \end{vmatrix} \cdot \begin{vmatrix} 1 & 0 & 1 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{vmatrix} \cdot \begin{vmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 1 & 0 & 1 \end{vmatrix} \cdot \begin{vmatrix} 1 & 0 & 2 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{vmatrix}$, qui

conduit du point extrême inférieur au point extrême supérieur de la

ligne limite extrême de droite, point à désigner, si $f' = \begin{pmatrix} 3, -1, -20 \\ 0, 0, 0 \end{pmatrix}$, par $\frac{3}{20} = \frac{a'}{A'}$. La ligne limite externe, conduisant du point $\frac{3}{20} = \frac{a_1}{A_1}$ vers la gauche, laquelle ne se termine pas à un point saillant, je me la figure, portée encore une fois, après qu'elle a franchi l'axe de symétrie, par lequel elle est coupée, conformément à la symétrie admise, et je désigne par R la substitution $\begin{vmatrix} -4 & 0 & 5 \\ 0 & -1 & 0 \\ -3 & 0 & 4 \end{vmatrix}$, composée des substitutions

$$\begin{vmatrix} 1 & 0 & 1 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{vmatrix} \cdot \begin{vmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 1 & 0 & 1 \end{vmatrix} \cdot \begin{vmatrix} -1 & 0 & 0 \\ 0 & -1 & 0 \\ -1 & 0 & 1 \end{vmatrix} \cdot \begin{vmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ -1 & 0 & 1 \end{vmatrix} \cdot \begin{vmatrix} 1 & 0 & -1 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{vmatrix},$$

qui mène du point $\frac{3}{20} = \frac{a_1}{A_1}$ à sa répétition $\frac{3}{20} = \frac{f_1(4, 0, 3)}{F_1(4, 0, -5)}$. Je me figure de même doublée la ligne limite ~~externe~~ conduisant du point $\frac{3}{20} = \frac{a'}{A'}$ vers la gauche, et je désigne par S la substitution $\begin{vmatrix} -2 & 1 & 0 \\ -3 & 2 & 0 \\ 0 & 0 & -1 \end{vmatrix}$

composée des substitutions $\begin{vmatrix} 1 & 0 & 0 \\ 1 & 1 & 0 \\ 0 & 0 & 1 \end{vmatrix} \cdot \begin{vmatrix} -1 & 1 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & -1 \end{vmatrix} \cdot \begin{vmatrix} 1 & 0 & 0 \\ -1 & 1 & 0 \\ 0 & 0 & 1 \end{vmatrix}$ et

qui conduit du point $\frac{3}{20} = \frac{a'}{A'}$ à sa répétition, le point $\frac{3}{20} = \frac{f'(2, 3, 0)}{F'(2, -1, 0)}$.

La figure quadruplée, conformément à la symétrie quadruple qui apparaît autour du point d'intersection des deux lignes limites encore restantes, est maintenant bornée seulement par quatre espèces de lignes, correspondant aux quatre substitutions P, Q, R et S, et que je désigne aussi par les lettres P, Q, R et S. Le long de la ligne limite de cette figure quadruplée, on arrive, en partant du point $\frac{12}{5} = \frac{a}{A}$, au point diamétralement opposé au moyen de la substitution

$$\begin{vmatrix} 29 & -5 & -15 \\ 60 & -11 & -30 \\ 36 & -6 & -19 \end{vmatrix}, \text{ qui répond au parcours aussi bien d'un groupe que}$$

de l'autre, de six fragments de limites qui y conduisent. Et par con-

$$\text{séquent, je remarque qu'en désignant } P.R.P^{-1} = \begin{vmatrix} -13 & 3 & 5 \\ -36 & 8 & 15 \\ -12 & 3 & 4 \end{vmatrix}$$

par T et $Q.S.Q^{-1} = \begin{vmatrix} -17 & 2 & 10 \\ -24 & 2 & 15 \\ -24 & 3 & 14 \end{vmatrix}$ par U, on aura $T.U = U.T$.

Si l'on se figure la ligne P prolongée par delà son point final, en produisant la figure précédente, conformément à la symétrie qui se reproduit à la ligne R, il faudra commencer par employer la substitution

$$\begin{vmatrix} -1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & -1 \end{vmatrix}, \text{ au moyen de laquelle la forme } \begin{pmatrix} 3, -4, -5 \\ 0, 0, 0 \end{pmatrix},$$

en vertu de cette symétrie, se change en elle-même. La substitution P^{-1} mène ensuite à la répétition prochaine du point $\frac{12}{5}$, de sorte que pour la forme f on obtient la substitution réciproque

$$\begin{vmatrix} -7 & 2 & 0 \\ -24 & 7 & 0 \\ 0 & 0 & -1 \end{vmatrix}, \text{ par laquelle elle se change en elle-même. Si l'on}$$

veut aussi franchir le dernier point final, on a encore, conformément à la symétrie qui se produit à la ligne Q, à ajouter la substitution

$$\begin{vmatrix} -1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & -1 \end{vmatrix}, \text{ de sorte qu'on obtient au total la substitution } \begin{vmatrix} 7 & 2 & 0 \\ 24 & 7 & 0 \\ 0 & 0 & 1 \end{vmatrix},$$

par laquelle f se transforme en elle-même : c'est exactement celle à laquelle aurait conduit l'emploi des prescriptions de II, b pour la forme $(12, 0, -1)$; je la désigne par V. Si l'on continue d'avancer sur la ligne P, il se manifeste nécessairement une concordance parfaite avec ce qui a été trouvé antérieurement, il part notamment du point

$\frac{12}{5} = \frac{f(7, 24, 0)}{F(7, -2, 0)}$, de même que du point $\frac{12}{5} = \frac{a}{A}$, une ligne Q, se dirigeant vers le même côté de la ligne P, et, après qu'on a employé la

substitution $\begin{vmatrix} -1 & 0 & 0 \\ 0 & -1 & 0 \\ 0 & 0 & 1 \end{vmatrix}$ correspondant à la transition de la ligne P,

semblablement une ligne Q se dirigeant vers l'autre côté. Tout ce qui a été dit de P s'applique par analogie à Q. Je désigne par W la substitution

$$Q. \begin{vmatrix} -1 & 0 & 0 \\ 0 & -1 & 0 \\ 0 & 0 & 1 \end{vmatrix}. Q^{-1}. \begin{vmatrix} -1 & 0 & 0 \\ 0 & -1 & 0 \\ 0 & 0 & 1 \end{vmatrix}.$$

On doit maintenant se figurer non-seulement les lignes primitives P et Q prolongées de part et d'autre, à l'infini, mais encore, pareillement, les lignes en nombre infini Q et P qui les croisent, puis celles

en nombre doublement infini qui croisent les dernières, etc., de telle sorte qu'on obtient un système de lignes semblable en quelque façon à un arbre ramifié en l'air et dans la terre. De plus, on doit se figurer qu'en un nombre infini de places séparées, tant sur le tronc que sur les branches, il passe toujours deux nouveaux rameaux dans des directions opposées, et qu'aucun rameau ne se trouve en contact avec un autre, excepté celui d'où il émane et ceux qui émanent de lui. La comparaison n'est défectueuse qu'en ce que chez moi tout doit être situé sur une surface correspondant à une des nappes hyperboloïdales et en ce que les rameaux devraient être de deux espèces différentes, en sorte que ceux d'une espèce ne produiraient que ceux de l'autre ; comme de toutes parts la ramification doit être poursuivie à l'infini, on peut faire coïncider chaque fragment d'une espèce avec chaque fragment de la même espèce au moyen de déformations appropriées. De chaque point du système à chaque autre il n'y a qu'un seul et même chemin. Ainsi, de chaque point, par exemple, de $\frac{12}{5} = \frac{a}{A}$ à chaque autre $\frac{12}{5}$ de ce système, il n'y a qu'une seule et même substitution, abstraction faite des changements admissibles de signes de colonnes tout entières, changements correspondant à la transition des lignes. Cette substitution est contenue dans l'expression générale $V^{\nu_0}.W^{\omega_0}.V^{\nu_1}.W^{\omega_1} \dots V^{\nu_x}.W^{\omega_x}$, que je veux désigner par E, et dans laquelle les exposants peuvent être des nombres entiers, positifs ou négatifs quelconques ; ν_0 et ω_x peuvent aussi être zéro. E n'est nullement l'expression la plus générale de toutes les substitutions, par lesquelles f se change en elle-même, car de tous les centres des fragments de l'espèce V se dirigent, vers les deux côtés de ces fragments, des lignes R, qui mènent à des centres analogues de fragments de l'espèce V appartenant à de semblables systèmes de lignes ; pareillement se dirigent des centres de tous les fragments de l'espèce W, vers les deux côtés de ces fragments, des lignes S, qui mènent à des centres analogues de fragments de l'espèce W appartenant à de semblables systèmes de lignes. Tous les nouveaux systèmes de lignes ainsi introduits ne sont en contact ni avec les systèmes primitifs ni avec eux-mêmes. Aussi les lignes R et S, partant à leur tour des points correspondants de ces systèmes, ne relient pas immédiatement

deux de ces systèmes de lignes, mais mènent vers de nouveaux et semblables systèmes de lignes, non cependant tous distincts, comme le montre l'étude du dodécagone, formé de lignes P, Q, R et S. On peut se faire une idée de toute cette configuration, si l'on se dit que le système de lignes considéré en premier lieu ne forme nulle part des polygones fermés, mais divise seulement la surface en un nombre infini de bandes s'étendant à l'infini. Dans le développement suivant naissent dans chacune de ces bandes de tels nouveaux systèmes de lignes, en nombre infini, nécessairement courbés convenablement, que je veux désigner par L_h , savoir un pour chaque fragment, ou V ou W de la délimitation de la bande, suivant que l'indice h est ou pair ou impair. Les bandes, en nombre infini, que chacun des nouveaux systèmes, de lignes L_h sépare d'avec la bande qu'on vient d'observer, doivent comme on voit facilement, être, à leur tour, traitées comme celle-ci, et ainsi de suite. Ce même rôle de la délimitation complète d'une pareille bande est joué par la ligne limite extrême d'un système de lignes L_h , laquelle, par la courbure admissible de toute la configuration, pourrait être changée en une délimitation semblable. Du nombre infini des systèmes de lignes, compris par cette limite, nous n'avons jusqu'ici considéré qu'un seul, le système primitif. Le suivant, d'un côté, coïncide avec un système appartenant, d'une manière analogue, à L_{h+1} ; et, de l'autre côté, le suivant coïncide avec un système appartenant, d'une manière analogue, à L_{h-1} .

Si l'on désigne par Γ et Δ les substitutions

$$\begin{vmatrix} -1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & -1 \end{vmatrix} \text{ et } \begin{vmatrix} -1 & 0 & 0 \\ 0 & -1 & 0 \\ 0 & 0 & 1 \end{vmatrix},$$

et si l'on remplace l'un des nombres t et u par 1, l'autre ainsi que γ et δ indifféremment par zéro ou 1, de la combinaison d'un nombre indéfini de substitutions E, Γ, Δ, T, U résulte une expression qui comprend en elle toutes les transformations de la forme $\begin{pmatrix} 12, & -1, & -5 \\ 0, & 0, & 0 \end{pmatrix}$ en elle-même et chacune seulement une fois, les substitutions E étant supposées

différentes de $\begin{vmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{vmatrix}$.

Pour traiter, de la manière la plus semblable possible, l'exemple de la *fig.* 5, je considère la forme $\begin{pmatrix} 1, & -2, & -15 \\ 0, & 0, & 0 \end{pmatrix}$; je désigne par Q la sub-

stitution $\begin{vmatrix} 2 & 1 & 0 \\ 1 & 1 & 0 \\ 0 & 0 & 1 \end{vmatrix}$, qui mène du point $\frac{1}{30}$ au point $\frac{2}{15}$, par R la substitution $\begin{vmatrix} -4 & 0 & 15 \\ 0 & -1 & 0 \\ -1 & 0 & 4 \end{vmatrix}$, réciproque, correspondant à la ligne limite

extrême placée à gauche et doublée, ligne menant du même point $\frac{1}{30}$ par-dessus le point $\frac{1}{15}$, en franchissant l'axe de symétrie supérieur de gauche à sa répétition $\frac{1}{30}$; je désigne par S la substitution

$\begin{vmatrix} -11 & 0 & 30 \\ 0 & -1 & 0 \\ -4 & 0 & 11 \end{vmatrix}$, correspondant à la ligne limite extrême placée à droite et

redoublée, tandis que la ligne désignée plus haut par P disparaît: alors

T devient égal à R, $U = Q.S.Q^{-1} = \begin{vmatrix} -21 & 20 & 60 \\ -10 & 9 & 30 \\ -4 & 4 & 11 \end{vmatrix}$ et l'on obtient de nouveau $T.U = U.T$. Je fais encore

$$Q. \begin{vmatrix} -1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & -1 \end{vmatrix}. Q^{-1}. \begin{vmatrix} -1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & -1 \end{vmatrix} = W;$$

alors, à la place du système précédent E de substitutions, se met simplement W^w ; à la place du système correspondant de lignes se met simplement une ligne allant de part et d'autre à l'infini. Du reste, il en est de même que plus haut, sauf que la substitution désignée par T disparaît.

Pour traiter la *fig.* 3 avec la plus grande analogie possible, je me l'imagine quadruplée autour du point $\frac{1}{4}$. Le long de la ligne limite inférieure et redoublée, la forme $\begin{pmatrix} 11, & -1, & -1 \\ 0, & 0, & 0 \end{pmatrix}$ se change alors en elle-même à l'aide de la substitution réciproque

$$\begin{vmatrix} 1 & 0 & 0 \\ 3 & 1 & 0 \\ 0 & 0 & 1 \end{vmatrix} \cdot \begin{vmatrix} 1 & 1 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{vmatrix} \cdot \begin{vmatrix} -1 & 1 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & -1 \end{vmatrix} \cdot \begin{vmatrix} 1 & -1 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{vmatrix} \cdot \begin{vmatrix} 1 & 0 & 0 \\ -3 & 1 & 0 \\ 0 & 0 & 1 \end{vmatrix} = \begin{vmatrix} -10 & 5 & 0 \\ -33 & 10 & 0 \\ 0 & 0 & -1 \end{vmatrix} = R.$$

La ligne limite primitive extrême de gauche mène d'un centre de

symétrie octuple, du point $\frac{11}{1}$, à un centre de symétrie sextuple, qui n'est pas un point $\frac{a}{A}$, sans qu'il en résulte plus de difficultés que dans les cas précédents. Il est un point de fissure $k = n = 0$ du champ de la forme $\begin{pmatrix} 3, & -2, & -2 \\ 1, & 1, & 0 \end{pmatrix}$, se changeant en elle-même à l'aide des substi-

tutions $\begin{vmatrix} 1 & 0 & 0 \\ 0 & -1 & 1 \\ 1 & -1 & 0 \end{vmatrix}$ et $\begin{vmatrix} -1 & 0 & 0 \\ -1 & 0 & 1 \\ -1 & 1 & 0 \end{vmatrix}$. A cette ligne limite correspond

la substitution $\begin{vmatrix} 1 & 0 & 0 \\ 2 & 1 & 0 \\ 2 & 0 & 1 \end{vmatrix}$, ou, si l'on ajoute encore la

substitution $\begin{vmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & -1 & 1 \end{vmatrix}$, pour obtenir la forme énoncée, celle-ci :

$\begin{vmatrix} 1 & 0 & 1 \\ 2 & 1 & 2 \\ 2 & -1 & 3 \end{vmatrix}$, que je désigne par Q. A la transition de la ligne limite

supérieure, dans la figure dessinée, au passage dans le deuxième

sextant correspond alors la substitution $\begin{vmatrix} -1 & 0 & 0 \\ -1 & 0 & 1 \\ -1 & 1 & 0 \end{vmatrix}$. Du point pri-

mitif $\frac{11}{1}$, on arrive maintenant à sa répétition diamétralement opposée, située dans la figure quadruplée, par deux voies auxquelles correspondent de nouveau les substitutions identiques l'une à l'autre T.U et U.T, si je mets de nouveau

$$R = T \quad \text{et} \quad Q \cdot \begin{vmatrix} -1 & 0 & 0 \\ -1 & 0 & 1 \\ -1 & 1 & 0 \end{vmatrix} \cdot Q^{-1} = \begin{vmatrix} -12 & 3 & 2 \\ -33 & 8 & 6 \\ -22 & 6 & 3 \end{vmatrix} = U.$$

Dans le système de lignes analogue à celui qui a été désigné jusqu'ici par E entre pareillement, comme dans l'exemple précédent, une simple et unique ligne, la ligne Q. Mais elle se dirige de chaque point $\frac{11}{1}$ en quatre sens et du point de fissure ainsi que de ses répétitions en trois sens; il en résulte qu'en désignant par Σ, Γ, Δ et Λ les sub-

stitutions $\begin{vmatrix} 1 & 0 & 0 \\ 0 & -1 & 1 \\ 1 & -1 & 0 \end{vmatrix}$, $\begin{vmatrix} -1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & -1 \end{vmatrix}$, $\begin{vmatrix} -1 & 0 & 0 \\ 0 & -1 & 0 \\ 0 & 0 & 1 \end{vmatrix}$ et $\begin{vmatrix} -1 & 0 & 0 \\ 0 & 0 & -1 \\ 0 & -1 & 0 \end{vmatrix}$, et

posant $\sigma = 1$ ou 2 , γ , δ et λ égaux à 0 ou 1 , il se produit évidemment une substitution E , par la combinaison de différentes substitutions formées comme $\Gamma^\gamma \Delta^\delta \Lambda^\lambda Q \Sigma^\sigma Q^{-1}$ et l'addition finale d'une autre substitution $\Gamma^\gamma \Delta^\delta \Lambda^\lambda$. Ici encore, comme dans le premier exemple, ce système de lignes a une propriété telle que, dût-on prolonger à l'infini la ramification, les différentes branches ne se rencontreront plus jamais. Dans les bandes infiniment nombreuses, qui résultent de la division de la surface par ce système, se retrouvent de nouveaux systèmes semblables, une ligne R partant de chaque point $\frac{1}{4}$ dans chacune des quatre bandes contiguës à ce point. Mais, comme le montre la figure quadruplée, étudiée ci-dessus, en face de toute la délimitation d'une pareille bande ne se trouvent que les lignes extrêmes d'un seul et même autre système de cette espèce. De même que plus haut, l'ensemble de ces lignes extrêmes peut être considéré aussi lui-même comme formant la délimitation d'une pareille bande. Une expression, qui donne toutes les transformations de la forme $\begin{pmatrix} 11, & -1, & -1 \\ 0, & 0, & 0 \end{pmatrix}$ en elle-même, est donc en tout cas celle-ci: $E_1 \cdot R \cdot E_2 \cdot R \cdot E_3 \cdot R \dots$, qui, par la propriété $T \cdot U = U \cdot T$, peut encore être amenée à un tel état qu'elle ne donne chaque substitution qu'une fois. Cela résulte de ce que, dans une substitution E , les derniers facteurs étant $\Lambda Q \Sigma Q^{-1}$, et donnant par conséquent ensemble U , on peut faire avancer devant eux le facteur R .

Avec la forme $\begin{pmatrix} 3, & -1, & -1 \\ 0, & 0, & 0 \end{pmatrix}$ (*fig. 4*), outre les substitutions Γ , Δ et Λ correspondant à son octuple symétrie, une seule substitution, réciproque, répondant à la ligne limite extrême de droite redoublée, la substitution

$$T = \begin{vmatrix} 1 & 0 & 0 \\ 1 & 1 & 0 \\ 0 & 0 & 1 \end{vmatrix} \cdot \begin{vmatrix} -1 & 1 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & -1 \end{vmatrix} \cdot \begin{vmatrix} 1 & 0 & 0 \\ -1 & 1 & 0 \\ 0 & 0 & 1 \end{vmatrix} = \begin{vmatrix} -2 & 1 & 0 \\ -3 & 2 & 0 \\ 0 & 0 & -1 \end{vmatrix}$$

suffit à donner une expression générale des transformations en elle-même de cette forme. Cette expression, qui résulte simplement de la combinaison de différentes substitutions $\Gamma^\gamma \Delta^\delta \Lambda^\lambda T$, peut être mise sous une telle forme qu'elle ne donne qu'une fois chaque substi-

tution. On y parvient par l'emploi de la proposition indiquée par la symétrie à douze parties au point $\frac{1}{3}$, savoir, que la sextuple ré-

pétition de la substitution $T.A = \begin{vmatrix} 2 & 0 & -1 \\ 3 & 0 & -2 \\ 0 & 1 & 0 \end{vmatrix}$ donne l'identité. Si

l'on voulait comparer le système de lignes actuel formé par les lignes T, dont quatre à la fois viennent aboutir à un même point, aux systèmes de lignes E étudiés dans les exemples posés jusqu'ici, on verrait donc ici les rameaux en contact en formant des hexagones.

Il y a évidemment, pour chaque forme $\begin{pmatrix} a, b, c \\ 0, 0, 0 \end{pmatrix}$, des substitutions pareilles à celles qui ont été désignées jusqu'ici par E. Un exemple d'une classe dans laquelle ne se produit aucune forme semblable est celui dont il est question (*fig. 2*). Je désigne maintenant par f

la forme $\begin{pmatrix} 2, -2, -1 \\ 0, 0, 1 \end{pmatrix}$ résultant, par la substitution $\begin{vmatrix} 2 & 0 & -1 \\ 0 & 1 & 0 \\ -1 & 0 & 1 \end{vmatrix}$,

de la forme $\begin{pmatrix} 1, -2, -2 \\ 1, 0, 1 \end{pmatrix}$ désignée plus haut par f et je cherche

ses transformations en elle-même. La substitution $R = \begin{vmatrix} -1 & 0 & 0 \\ -1 & 1 & 0 \\ 0 & 0 & -1 \end{vmatrix}$

mène du point $\frac{2}{2} = \frac{f(1, 0, 0)}{F(1, 0, 0)}$ à un point symétrique $\frac{2}{2} = \frac{f(1, 1, 0)}{F(1, 0, 0)}$,

situé au delà d'un axe de symétrie non dessiné, et, pour aller de celui-ci au point $\frac{2}{2}$ suivant, sur le contour dessiné de la figure, on

emploie la substitution

$$S = \begin{vmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 1 & 0 & 1 \end{vmatrix} \cdot \begin{vmatrix} -1 & -1 & 2 \\ 0 & -1 & 0 \\ 0 & -1 & 1 \end{vmatrix} \cdot \begin{vmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ -1 & 0 & 1 \end{vmatrix} = \begin{vmatrix} -3 & -1 & 2 \\ 0 & -1 & 0 \\ -4 & -2 & 3 \end{vmatrix}.$$

La figure montre que la substitution R.S, après avoir été employée trois fois, donne l'identité. Si l'on désigne par Δ et Φ les substitu-

tions $\begin{vmatrix} -1 & 0 & 0 \\ 0 & -1 & 0 \\ 0 & 0 & 1 \end{vmatrix}$ et $\begin{vmatrix} -1 & -1 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & -1 \end{vmatrix}$, qui correspondent à la transi-

tion des deux espèces d'axes de symétrie dessinés, si l'on met δ et φ égaux à 0 ou 1, et si l'on fait toutes les dispositions possibles des

substitutions R, S, Δ^3, Φ^9 , on obtient de la sorte une expression pour toutes les transformations de la forme f en elle-même.

Les relations

$$R\Delta = \Delta R, \quad S\Phi = \Phi S, \quad \Delta\Phi = \Phi\Delta, \quad R^3 = S^3 = \Delta^3 = \Phi^3 = \begin{vmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{vmatrix}$$

montrent qu'on peut avancer la substitution Δ de droite à gauche toujours assez loin pour qu'elle vienne s'arrêter immédiatement après S ; de même la substitution Φ , pour qu'elle vienne s'arrêter immédiatement après R , de sorte que l'expression entière ne contient plus, outre le facteur initial $\Delta^3\Phi^9$, que des puissances positives de

$$R\Phi = \begin{vmatrix} 1 & 1 & 0 \\ 1 & 2 & 0 \\ 0 & 0 & 1 \end{vmatrix}, \quad S\Delta = \begin{vmatrix} 3 & 1 & 2 \\ 0 & 1 & 0 \\ 4 & 2 & 3 \end{vmatrix},$$

parmi lesquelles peuvent encore être réparties des connexions de R et S seules, mais seulement R avant $S\Delta$, puis S et RS avant $R\Phi$. En effet, à cause de la rela-

$$\text{tion } (RS)^3 = \begin{vmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{vmatrix}, \text{ toutes les connexions de } R \text{ et } S \text{ peuvent}$$

d'ailleurs être supprimées, sauf RS, SR, RSR .

Ces exemples suffiront pour montrer la variété qui se manifeste, même dans les formes les plus simples, variété qui, pour le moment, ne fait apparaître comme impossible l'établissement d'une formule valable pour toutes les formes à la fois, ce qui ne veut pas dire que je nie la possibilité d'une certaine classification, qui sera cependant beaucoup plus compliquée que la classification analogue pour les formes positives.

Dans les articles déjà cités de M. Hermite, qui découvrit la réduction continue, se trouve énoncée dans la proposition I, p. 312 (t. 47 du *Journal de Crelle*), que, si la forme f se change en elle-même, par

$$\text{la substitution } S = \begin{vmatrix} \alpha & \beta & \gamma \\ \alpha_1 & \beta_1 & \gamma_1 \\ \alpha_2 & \beta_2 & \gamma_2 \end{vmatrix}, \text{ l'équation } \begin{vmatrix} \alpha-t & \beta & \gamma \\ \alpha_1 & \beta_1-t & \gamma_1 \\ \alpha_2 & \beta_2 & \gamma_2-t \end{vmatrix} = 0$$

possède une racine $t_1 = 1$, ou, ce qui revient au même, que la somme des éléments de la diagonale principale de la matrice donnée $\alpha + \beta_1 + \gamma_2$ est égale à la somme analogue formée d'éléments de la matrice de la substitution adjointe, ou qu'on a $t = 1$ dans un des trois systèmes de

nombre λ, μ, ν, t , qui donnent identiquement

$$\lambda x + \mu y + \nu z = t(\lambda x' + \mu y' + \nu z').$$

(voir Cayley, t. 50, p. 291, et Bachmann, t. 76, p. 335 du même journal). En effet, si f se transforme par la substitution S en f' , identique à f , il y a, dans les champs de F et F' , deux points, que je distingue par les indices 1 et 2, et qui sont déterminés par des valeurs

$$\begin{aligned}\Xi_1 = \Xi_2 = \alpha \Xi'_1 + \beta H'_1 + \gamma Z'_1, \quad H_1 = H_2 = \alpha_1 \Xi'_1 + \beta_1 H'_1 + \gamma_1 Z'_1, \\ Z_1 = Z_2 = \alpha_2 \Xi'_1 + \beta_2 H'_1 + \gamma_2 Z'_1,\end{aligned}$$

à rapports rationnels entre elles. L'équation de la ligne $\begin{vmatrix} \Xi' & H' & Z' \\ \Xi'_1 & H'_1 & Z'_1 \\ \Xi'_2 & H'_2 & Z'_2 \end{vmatrix} = 0$

réunissant ces deux points, je la dénote par $\lambda \Xi' + \mu H' + \nu Z' = 0$. Soient F_p et F'_p des formes provenant de F et F' par une substitution dont

la dernière colonne est $\begin{vmatrix} \lambda \\ \mu \\ \nu \end{vmatrix}$; alors $\begin{vmatrix} 0 \\ 0 \\ 1 \end{vmatrix}$ est la dernière colonne de la

substitution qui correspond à ladite ligne, et par laquelle $\mathcal{F}_p$, formée au moyen de Ξ_2, H_2, Z_2 , se change en $\mathcal{F}'_p$ formée au moyen de Ξ'_2, H'_2, Z'_2 , par laquelle donc F_p se change en F'_p . Donc 001 est la dernière ligne de la substitution par laquelle f_p se change en f'_p , et l'on a

$$\lambda x + \mu y + \nu z = z_p = z'_p = \lambda x' + \mu y' + \nu z'.$$

Quant aux deux autres valeurs admissibles t_2 et t_3 de t , réciproques l'une à l'autre, qui font identiquement $\lambda x + \mu y + \nu z = t(\lambda x' + \mu y' + \nu z')$, il faut, comme le dit la proposition II de M. Hermite, si

$S^n = \begin{vmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{vmatrix}$ et que n soit le plus petit nombre de cette propriété,

qu'elles soient des racines primitives $n^{\text{ièmes}}$ de l'unité. Toutefois 6 peut être aussi la valeur de n , comme le montre mon développement à la fin de IV b, et l'exemple mentionné de la forme $\begin{pmatrix} 3, -1, -1 \\ 0 & 0, & 0 \end{pmatrix}$.

On peut constater, même sans ces développements, que n n'a pas d'autres valeurs possibles que 1, 2, 3, 4, 6. En effet, tandis que, d'après l'équation $t^2 + (1 - \alpha - \beta_1 - \gamma_2)t + 1 = 0$, valable pour t_2 et t_3 ,

les valeurs $-1, 0, 1, 2, 3$ de $\alpha + \beta_1 + \gamma_2$ donnent pour t_2 et t_3 des racines primitives respectivement $2^{\text{es}}, 3^{\text{es}}, 4^{\text{es}}, 6^{\text{es}}, 1^{\text{re}}$ de l'unité, dans toutes les autres valeurs entières de cette somme, t_2 et t_3 deviennent réels et aussi différents de ± 1 . M. Hermite avait évidemment l'intention d'exclure de ses autres recherches ces substitutions, appartenant aux différents nombres n , qui sont analogues aux simples racines de l'unité, se produisant dans les transformations semblables des formes divisibles en facteurs linéaires, ou, ce qui revient au même, dans les unités complexes. Car, dans le cas $n = 2$, ou $t_2 = t_3 = -1$, les équations (5) (page 309, t. 47 du *Journal de Crelle*) ne seraient pas les plus générales satisfaisant à l'équation (4), parce qu'alors est égal à zéro le déterminant du système de coefficients, au moyen duquel, d'après ma désignation, $x + x', y + y', z + z'$ sont exprimés par x', y', z' ; alors ne serait non plus valable la proposition (page 324) que T.U ne peut pas être égal à U.T, sans que les deux substitutions soient des puissances d'une seule et même substitution. En ce cas on ne doit pas appliquer non plus la formule donnée par M. Cayley, au passage indiqué (page 293), attendu qu'alors les coefficients qui sont désignés là par B et C peuvent aussi rester différents de zéro. Pour $n = 3, 4, 6$ on peut employer les formules générales, nonobstant que les valeurs de t_2 et t_3 sont alors complexes, et que par conséquent, d'après les expressions données pour elles, page 318, par M. Hermite, la quantité désignée là par γ , et que je désignerais par $-F(\lambda, \mu, \nu)$, ne peut alors pas être positive. Au fond les formules de M. Hermite reviennent aux répétitions de la valeur numérique de a dans des territoires de coefficients négatifs A, tels que je les désigne, répétitions considérées plus haut, et que l'on ne doit pas modifier essentiellement, même pour des formes non réduites. Elles donnent, de la manière la plus facile, les substitutions semblables prises une à une; c'est seulement pour constater leur ensemble et leur dépendance mutuelle que la considération des territoires de coefficients positifs A et a est utile, à cause de leur propriété de s'exclure mutuellement.

La quantité $\alpha + \beta_1 + \gamma_2$ a la propriété, quand $\begin{vmatrix} \alpha & \beta & \gamma \\ \alpha_1 & \beta_1 & \gamma_1 \\ \alpha_2 & \beta_2 & \gamma_2 \end{vmatrix}$ est une substitution composée, et que l'on modifie l'ordre de ses composantes,

de rester invariable, attendu qu'elle devient $\sum_p \sum_\sigma p_{p\sigma} \cdot q_{p\sigma}$, quand $p_{p\sigma}$ et $q_{p\sigma}$ sont les éléments des composantes. On peut donc, en quelque sorte, regarder cette grandeur ou une de ses fonctions appropriées, par exemple, $\frac{\alpha + \beta_1 + \gamma_2 - 1}{2}$ ou le module ou la portion réelle du logarithme de l'une des deux racines de l'équation

$$t^2 - (\alpha + \beta_1 + \gamma_2 - 1)t + 1 = 0,$$

comme échelle du rang d'une substitution, ce qui devient particulièrement simple et naturel, lors de l'exclusion des substitutions, plusieurs fois mentionnées et devant être regardées comme des exceptions, pour lesquelles $n = 2, 3, 4, 6$. D'après la proposition précitée, il y a, chaque fois, après le choix de cette échelle, égalité ou une autre relation simple, entre le rang d'une substitution et celui de son inverse ou de son adjointe. Dans ce que l'on vient de dire, on a évidemment déjà résolu le problème d'exprimer toutes les substitutions, par lesquelles f se transforme en elle-même, au moyen de quelques substitutions indépendantes les unes des autres et ayant le rang absolument le plus petit, et ce problème équivant, par analogie, au problème de ramener toutes les substitutions semblables, des formes binaires, à la substitution, répondant à une simple période ou à la plus petite solution de l'équation de Pell.

f. — Des formes par lesquelles on peut rationnellement représenter le nombre zéro.

Dans ces formes, d'où je commence par exclure le cas $I = 0$, continue à subsister la définition de formes réduites et la division en champs et en territoires. Seulement si, parmi les nombres a, b, c, d , zéro lui-même apparaît, s'ajoutent aux valeurs de λ mentionnées en IV b et aux substitutions correspondantes des substitutions ultérieures, vu qu'alors, des nombres g, h, k, l, m, n , quatre peuvent devenir égaux à zéro et les champs de pareilles formes réduites peuvent s'étendre à l'infini. Je ne veux plus faire ici de développements particuliers, mais seulement appliquer immédiatement les principes généraux parfaitement suffisants, d'abord à l'exemple de la forme

$\begin{pmatrix} 1, & -1, & -1 \\ 0, & 0, & 0 \end{pmatrix} = f$, traité dans les *OEuvres posthumes* de Gauss (t. II, p. 311), seul exemple d'une forme ternaire indéfinie traité jusqu'à présent, qui soit parvenu à ma connaissance. Cet exemple est si simple, que le dessin à tracer, d'après le principe que j'ai utilisé antérieurement, se réduirait à la moitié d'un seul champ. Car le

champ de la forme f , $= \begin{vmatrix} 0 & 0 & 1 & -1 \\ -1 & 0 & 1 \\ -1 & 0 \\ 0 \end{vmatrix}$, résultant de f au moyen de

la substitution $\begin{vmatrix} 1 & 0 & 0 & -1 \\ -1 & 0 & 1 & 0 \\ 0 & -1 & 0 & 1 \end{vmatrix}$, est symétrique à lui-même, corres-

pondant à la substitution $\begin{vmatrix} -1 & 0 & 0 & 1 \\ -1 & 0 & 1 & 0 \\ -1 & 1 & 0 & 0 \end{vmatrix}$ et à l'axe $\eta = \zeta$, et il est

séparé par chacune de ses trois lignes limites, dont il ne faut considérer qu'une dans son parcours entier, et une autre, dans la moitié de son parcours, d'avec un champ qui lui est symétrique; savoir, il est séparé par le fragment, s'étendant de $\zeta = 0$ jusqu'à $\zeta = +\infty$, de la ligne $\eta = 0$, d'avec le champ de la forme qui concorde numériquement avec elle et qui en résulte au moyen de la substitution

$\begin{vmatrix} -1 & 0 & 0 & 1 \\ 0 & -1 & 0 & 1 \\ -2 & 0 & 1 & 1 \end{vmatrix}$. Le long du fragment, qui s'étend de $\eta = 0$,

$\zeta = +\infty$ jusqu'à $\eta = \zeta = \sqrt{\frac{1}{2}}$, de la ligne $2\eta\zeta = 1$, on a $\eta + \zeta = \xi$ d'après (11) et attendu que ξ, η, ζ sont positifs. Si, pour répondre au passage de la ligne limite, on met $\eta + \zeta = \xi + \vartheta$, on a

$$2\eta\zeta = 1 + 2\xi\vartheta + \vartheta^2, \quad h_1 = 2\eta(\xi - \eta) - 1 = 2(\xi - \eta)\vartheta + \vartheta^2,$$

$$l_1 = -\vartheta^2, \quad m_1 = 2(\xi - \zeta)\vartheta + \vartheta^2.$$

Pour traiter ce cas, qui n'eût pas été possible avec les formes considérées antérieurement, il faut d'abord, conformément à la prescription de

III a, appliquer la substitution $\begin{vmatrix} 1 & -1 & 0 & 0 \\ 0 & -1 & 0 & 1 \\ 0 & 0 & -1 & 1 \end{vmatrix}$, dans laquelle seulement l'ordre des colonnes est arbitraire et qui, à la place de la somme

$-g_1, -h_1, -k_1, -l_1, -m_1, -n_1$, amène une somme diminuée de h_1 , qui est le plus grand des nombres $g_1, h_1, k_1, l_1, m_1, n_1$, ensuite, parce que le coefficient remplaçant k_1 devient seul positif, il faut appliquer la substitution

$\begin{vmatrix} 0 & 1 & 0 & -1 \\ -1 & 0 & 1 & 0 \\ 0 & 0 & 1 & -1 \end{vmatrix}$, par conséquent il faut appliquer en tout à f_1 la

substitution $\begin{vmatrix} 1 & 1 & -1 & -1 \\ 1 & 0 & -1 & 0 \\ 0 & 0 & -1 & 1 \end{vmatrix}$. Comme dans la forme résultante, que je

nomme f_2 , on obtient $g_2 = -\delta^2$, $h_2 = 2(\xi - \zeta)\delta$, $k_2 = -2(\xi - \eta)\delta$, d'où l'on peut reconnaître que son champ dégénère en la ligne $2\eta\zeta = 1$, on aurait encore à employer des substitutions ultérieures, et tout d'abord une qui change h_2 en $-h_2$. Mais on prévoit que, dans le cas présent, qui est celui d'une valeur positive de δ , si l'on a mis en

avant la substitution $\begin{vmatrix} 1 & 0 & 0 \\ 0 & -1 & 0 \\ 0 & 0 & -1 \end{vmatrix}$, on peut employer la même substitution

pour une réduction ultérieure qui, dans le cas d'une valeur négative de δ , aurait changé la forme f_2 en la forme f_1 . Nous ferons donc

usage de l'inverse de la précitée, $\begin{vmatrix} 0 & 1 & -1 & 0 \\ 1 & -1 & 0 & 0 \\ 0 & 0 & -1 & 1 \end{vmatrix}$, qui, attendu que la

forme $f_2 = \begin{pmatrix} -1, 0, 0 \\ -1, 0, 0 \end{pmatrix}$, n'a pas été changée numériquement par la

substitution $\begin{vmatrix} 1 & 0 & 0 \\ 0 & -1 & 0 \\ 0 & 0 & -1 \end{vmatrix}$, produit une forme concordant numériquement

avec la forme f_1 , et résultant d'elle, comme on voit, par la substitution

$\begin{vmatrix} -1 & 2 & -2 \\ 0 & 1 & -2 \\ 0 & 0 & -1 \end{vmatrix}$. A l'aide des trois substitutions, réciproques à

elles-mêmes, qu'on a trouvées, on peut former toutes les substitutions, par lesquelles f_1 se change en elle-même. Si on les fait précéder

de la substitution $\begin{vmatrix} 1 & 0 & 0 \\ -1 & 0 & 1 \\ 0 & -1 & 0 \end{vmatrix}$ et suivre par l'inverse, on obtient les

substitutions

$$S = \begin{vmatrix} -1 & 0 & 0 \\ 0 & 0 & -1 \\ 0 & -1 & 0 \end{vmatrix}, \quad T = \begin{vmatrix} -1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & -1 \end{vmatrix}, \quad U = \begin{vmatrix} -3 & -2 & -2 \\ 2 & 1 & 2 \\ 2 & 2 & 1 \end{vmatrix},$$

au moyen desquelles on peut former toutes les substitutions qui changent en elle-même la forme $\begin{pmatrix} 1, & -1, & -1 \\ 0, & 0, & 0 \end{pmatrix}$. L'expression générale de ces substitutions peut se formuler d'une manière telle qu'elle ne donne chacune qu'une fois, au moyen des relations $S.U = U.S$, $STST = TSTS$, valables pour les substitutions S, T, U , ainsi que pour les trois nommées en premier lieu; ces relations résultent de la symétrie quadruple et octuple, qui se produit à deux points d'intersection des trois axes de symétrie considérés, points situés dans le fini. La symétrie, qui se produit au troisième point situé dans l'infini, a lieu entre une infinité de parties. Il n'est pas difficile de voir que le système de coefficients donné par Gauss

$$W = \begin{vmatrix} \frac{1}{2}(\alpha^2 + \beta^2 + \gamma^2 + \delta^2) & \frac{1}{2}(\alpha^2 + \beta^2 - \gamma^2 - \delta^2) & \alpha\gamma + \beta\delta \\ \frac{1}{2}(\alpha^2 - \beta^2 + \gamma^2 - \delta^2) & \frac{1}{2}(\alpha^2 - \beta^2 - \gamma^2 + \delta^2) & \alpha\gamma - \beta\delta \\ \alpha\beta + \gamma\delta & \alpha\beta - \gamma\delta & \alpha\delta + \beta\gamma \end{vmatrix},$$

dans lequel $\begin{vmatrix} \alpha & \beta \\ \gamma & \delta \end{vmatrix} = 1$, et dans lequel ou deux des nombres $\alpha, \beta, \gamma, \delta$ sont entiers et pairs, les deux autres entiers et impairs, ou tous les quatre sont des multiples impairs de $\sqrt{\frac{1}{2}}$, donne en réalité les transformations de la forme $\begin{pmatrix} 1, & -1, & -1 \\ 0, & 0, & 0 \end{pmatrix}$ en elle-même, à l'exclusion, ce qui va sans dire, de celles qui ont un premier coefficient négatif, et que l'on peut aisément déduire des autres. Car d'abord, si l'on remplace

les nombres $\begin{vmatrix} \alpha & \beta \\ \gamma & \delta \end{vmatrix}$ par les nombres $\begin{vmatrix} 1 & \pm 2 \\ 0 & 1 \end{vmatrix}$ ou $\begin{vmatrix} \sqrt{\frac{1}{2}} & \mp \sqrt{\frac{1}{2}} \\ \pm \sqrt{\frac{1}{2}} & \sqrt{\frac{1}{2}} \end{vmatrix}$,

le système de coefficients coïncide avec celui de la substitution

$$(STSU)^{\pm 1} = \begin{vmatrix} 3 & 2 & \pm 2 \\ -2 & -1 & \mp 2 \\ \pm 2 & \pm 2 & 1 \end{vmatrix} \text{ ou } (ST)^{\pm 1} = \begin{vmatrix} 1 & 0 & 0 \\ 0 & 0 & \pm 1 \\ 0 & \mp 1 & 0 \end{vmatrix}, \text{ tandis que,}$$

si l'on fait permuter α avec δ , β avec γ , de W résulte le système $T.W.T$. A cette relation on pourrait joindre une série de relations semblables, basées sur des changements de signes et sur d'autres permutations entre $\alpha, \beta, \gamma, \delta$, comme par exemple si, dans $\begin{vmatrix} \alpha & \beta \\ \gamma & \delta \end{vmatrix}$,

on fait permuter les deux lignes horizontales, et si, dans l'une d'elles,

on change les signes, W se change en $WSTST = W \begin{vmatrix} 1 & 0 & 0 \\ 0 & -1 & 0 \\ 0 & 0 & -1 \end{vmatrix}$.

D'un côté, toutes les séries de facteurs S, T, U peuvent maintenant s'exprimer par les expressions $STSU, ST$ qu'on fera précéder ou suivre à la fois de T , abstraction faite d'un facteur T , qui peut être surabondant au commencement ou à la fin. De l'autre côté, tous les systèmes $\begin{vmatrix} \alpha & \beta \\ \gamma & \delta \end{vmatrix}$ de la première et de la seconde espèce peuvent se

composer à l'aide des systèmes simples $\begin{vmatrix} 1 & \pm 2 \\ 0 & 1 \end{vmatrix}$ et $\begin{vmatrix} \sqrt{\frac{1}{2}} & \pm \sqrt{\frac{1}{2}} \\ \pm \sqrt{\frac{1}{2}} & \sqrt{\frac{1}{2}} \end{vmatrix}$,

de la manière qu'a lieu la composition des substitutions, de même que la composition de tels systèmes quelconques donnant de nouveau des systèmes semblables; et il est aisé de voir que, lorsqu'un système $\begin{vmatrix} \alpha'' & \beta'' \\ \gamma'' & \delta'' \end{vmatrix}$ est égal à $\begin{vmatrix} \alpha & \beta \\ \gamma & \delta \end{vmatrix} \cdot \begin{vmatrix} \alpha' & \beta' \\ \gamma' & \delta' \end{vmatrix}$, le système correspondant W'' est égal à $W' \cdot W$.

La *fig. 6* suffira pour expliquer un exemple ultérieur, la forme $\begin{pmatrix} 1, & -1, & -47 \\ 0, & 0, & 0 \end{pmatrix}$. Les frontières des territoires de coefficients positifs A y sont marquées par des lignes pleines renforcées, au lieu de l'être comme auparavant par des lignes serpentantes. Parmi les lignes limites externes, il n'y a que six axes de symétries. Au delà des cinq autres, la continuation est obtenue par la demi-révolution de celles-ci autour des centres de symétries, qui y sont marqués par des croix; le centre situé dans le territoire de $a = 2$ rentre dans ce qui a été dit III, *b*, 6, les autres sont eux-mêmes points de croisement.

Dans le cas $I = 0$, il résulte de $h = k = 0$, à cause de $aA = 0$, que, ou bien $a = \xi = a = h = k = 0$, et par conséquent $f = (b, g, c)$, ou bien que $A = \Xi = A = H = K = 0$. Si, dans le deuxième cas, r désigne le plus grand commun diviseur de h et k , q celui de b et c , il en résulte que

$$b = \pm q \frac{k^2}{r}, \quad c = \pm q \frac{h^2}{r}, \quad g = \pm q \frac{hk}{r},$$

$$f = ax^2 + 2rx \left(\frac{k}{r} y + \frac{h}{r} z \right) \pm q \left(\frac{k}{r} y + \frac{h}{r} z \right)^2.$$

Dans les deux cas, la forme indéfinie réduite f devient donc une forme binaire, soit indéfinie, soit négative, de sorte que f même doit être nommée une forme indéfinie ou une forme non positive. Car, si l'on désigne f par (b, g, c) , si (b, g, c) était une forme binaire positive, les conditions de réduction ne pourraient être remplies par aucune forme positive (b, g, c) .

Je ne puis finir sans rendre grâce à l'extrême obligeance de M. Hermite, qui a bien voulu prendre la peine de revoir avec le plus grand soin tout le travail précédent [*].

[*] ERRATA. — Page 21, note, *au lieu de XXX, lisez XXXIII*. — Page 60, ligne 11, *au lieu de : 4, 8, 12, lisez : 4, 8, 24*.